

THE COMPLETE
TECHNICAL PAPER PROCEEDINGS
FROM:



A COMPARATIVE ANALYSIS OF IP STREAMING VIDEO VERSUS MPEG VIDEO OVER CABLE

¹
Prashant P. Gandhi, Ph.D.

Cisco Systems, Inc.

Abstract

Cable networks have been designed to carry broadcast digital video services for consumer entertainment. Recent popularity of cable high-speed data service indicates the possible use of the cable DOCSIS infrastructure to deliver streaming media services. This paper presents a high-level comparison of video delivery over MPEG and over IP transport mechanisms. Issues such as broadcast and point-to-point transmissions, quality of service (QoS) and set-top box support are discussed.

1. INTRODUCTION

The MPEG-2 standard for compression and transport of entertainment-quality video has been critical for offering digital television (DTV) services to cable consumers. The broadcast nature of MPEG transport allows its efficient delivery over a shared cable infrastructure. The quality and convenience of digital programming has made the DTV service quite popular in the cable community. Digital set-top boxes (STBs) capable of decoding MPEG-2 video are being deployed in growing numbers in consumer homes.

The two-way capable hybrid fiber coax (HFC) cable plant has also allowed offering of DOCSIS data service to consumers. Cable modem service for high-speed internet access has been making significant in-roads in internet-hungry homes. By

the end of 2000, nearly 5.5 million US and Canadian homes subscribed to cable modem service [1]. Of these, 3 million subscribers were added in 2000 (with 1 million added in the 4th quarter alone). Appetite for broadband cable data service is also fueled by telecommuters that require high-speed access to corporate intranets and whose monthly subscription fees are subsidized by their employers.

Broadband has enabled the delivery of internet-based video (a la streaming media) to the PC. Growing number of websites are now providing streaming video at broadband bitrates, typically in the range 128 kb/s – 1 Mb/s. These are both broadcast (multicast) programs, such as television and radio channels, and on-demand programs such as news, sports, music, music videos, international, movie trailers and independent films. Also, Internet video is typically encoded in proprietary formats, such as Real Networks, Microsoft Windows Media Technology (WMT) and Apple QuickTime (QT).

With nearly 70 million cable-ready homes, the growth of digital TV and data services are expected to remain in heavy demand for the next several years. Clearly, these two services have traditionally addressed two different market segments. The DTV service provides television entertainment in the family room whereas high-speed data (HSD) service provides internet access to PC in the study.

¹ Prashant Gandhi (gandhi@cisco.com) is a Technical Leader in the Service Provider Line of Business at Cisco Systems.

This differentiating line between DTV and HSD services have now started to blur with the introduction of interactive television (iTV) STBs. These STBs, running iTV middleware software (e.g. from Liberate or Microsoft), integrate the television experience with internet-based broadband interactivity. This broadband interactivity is achieved by the integration of a DOCSIS cable modem with a digital STB.

Internet streaming media, as it stands today, is mostly viewed at little or no cost to consumers (other than the cost of HSD service). With broadband interactivity to the TV, there is an opportunity to provide IP streaming media services over DOCSIS. These services may be provided as part of a premium package or as part of additional subscription and/or pay-per-view fees. They can be delivered to a PC or a TV.

Niche content, generally not available on the television line-up, can be offered over cable's IP/DOCSIS network. Dynamic content, such as news and sports, international content, children's programming, etc. are examples of content that consumers may pay for to watch on demand. The value is in the flexibility of on-demand service in terms of convenience, choice and control. Portability of content can also be a value add – consumer should be able to watch a program on PC and transfer it to TV at a moment's notice.

This paper presents a comparison between MPEG and IP video and discusses some of the issues and challenges in deploying IP video to the TV.

2. CHARACTERISTICS OF MPEG AND IP VIDEO

MPEG Video [2,3]

In MPEG, the source of video is typically an MPEG-2 compressed video bitstream (in the

form of packetized elementary stream or PES) that is sent to an MPEG transport stream (TS) multiplexor. The multiplexor takes one or more elementary streams and produces a single- or multi-program transport stream.

Each PES is packetized into 188-byte packets, with 184-byte payload and 4-byte header. Video, audio and private data streams are individually packetized. Video and audio streams are also time synchronized to help guide the decoder assembled a synchronized audio/video presentation (a la lip synchronization).

Each packetized stream is marked with a packet ID (PID). Audio/video/data PIDs of a specific program are sent separately as part of program specific information (PSI) tables. PSI tables include: program allocation table (PAT) with PID = 0 to provide map information of all transport stream programs, program map table (PMT) indicating PIDs for each program within the TS, conditional access table (CAT) with PID = 1 for PIDs of entitlement messages and tables that carry private data.

The MPEG TS is potentially encrypted and then protected via forward error correction (FEC) as per ITU J83 Annex B FEC. Quadrature amplitude modulation (QAM) is then applied to the FEC stream with 6 bits/symbol (QAM-64) or 8 bits/symbol (QAM-256).

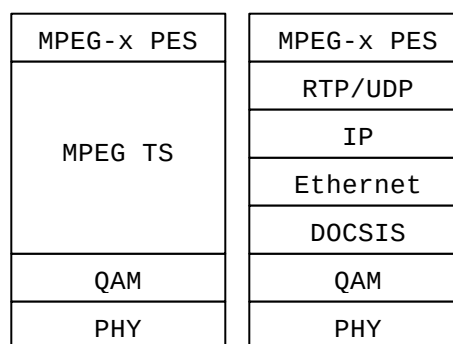


Figure 1: MPEG and IP Video Stack

As indicated earlier, MPEG TS packets do not carry any source or destination information. For a point-to-point video session (e.g. for on-demand video), PID information and entitlement information (for an encrypted stream) must be sent to the client. Also, for a point-to-point session, a control session needs to be established between client and video server. This is typically done via proprietary protocol or via MPEG's digital storage media – command and control (DSM-CC) protocol. Due to the complexity and interoperability issues with DSM-CC, there has been some movement to adopt the real-time streaming protocol (RTSP) for session control.

IP Video

With IP, both point-to-point (unicast) and point-to-multipoint (multicast) transmission is possible. This allows on-demand and broadcast video services to be offered on a single infrastructure. In addition, the IP infrastructure can be shared among multiple services, including data, telephony and video.

The video source for Internet protocol (IP) based transport can still be MPEG-1/2/4 PES. The PES stream is converted to a stream of IP packets which are then transported via the unreliable user datagram protocol (UDP). Each UDP/IP packetized video or audio stream is appended with real-time transport protocol (RTP) header to allow stream synchronization, time stamps and sequence numbering [4].

Typically, to avoid fragmentation, each RTP/UDP/IP packet size is less than one Ethernet frame payload (~1470 bytes). RTP/UDP/IP header overhead is at least 40 bytes (12 bytes for RTP, 8 bytes for UDP and minimum 20 bytes for IP). Clearly, the IP packet overhead is much larger than that for MPEG TS packets. However, for larger video packets (~1470 bytes), this overhead is not as significant.

Also, RTP packet header compression may be employed to reduce the 40-byte overhead to only several bytes [5]. Standardized mechanism of packetizing MPEG-1/2 video over RTP is specified in IETF RFC 2250 [6].

Real-time streaming protocol (RTSP) is employed to control a video session [7]. RTSP opens a separate connection (typically TCP/IP) between the client and video server. This is shown in Figure 2. RTSP allows session initiation, VCR-like session control such as play, pause, stop, etc., and session termination. RTP session parameters are communicated using RTSP. An optional communication channel via real-time control protocol (RTCP) can also be negotiated to allow the client to communicate additional information (e.g. packet loss) to the server. RTSP is also used to communicate multicast program information to the client.

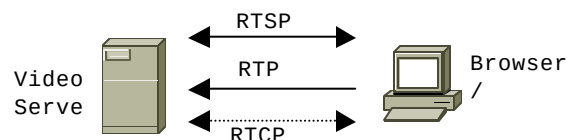


Figure 2: Client-server interactions for IP Video

Though it is important to support standards-based video for coding, transport and storage, today's Internet streaming media is dominated by proprietary systems of Real Networks, Microsoft and Apple. This is shown in Figure 3. Real Networks uses proprietary transport called Real Data Transport (RDT), proprietary G2 codec and proprietary file format (e.g. .rm, .ra). Microsoft WMT uses the proprietary control and transport protocol MMS and proprietary ASF file format. WMT does support standard MPEG-1/2 codecs as well as its own proprietary audio/video codecs. Apple QT uses the standards-based RTSP and RTP for control and transport and its file format has been adopted (with some variation) for storing MPEG-4 video.

Vendor	Product	Platform	Protocols	Codecs	File Formats
Real Network	Real	NT, Unix	RTSP, RDP	G2 (Proprietary)	RM
Microsoft	WMT	NT/W2K	MMS	Proprietary and MPEG-1/2	ASF
Apple	QuickTime	MacOS, Linux	RTSP, RTP	H.261, Sorenson (proprietary)	QT

Figure 3: IP Streaming Media Formats

However, its primary codec is proprietary (supplied by Sorenson).

Proprietary nature of IP video is not as critical on the PC as it is on the STB. STB's resources utilization (processing power, run-time memory, codesize, etc), player stability and frequent software upgrades are issues that need to be addressed for a scalable deployment of IP video to TVs.

Streaming Media Session Interactions

The following illustrates the interactions for establishing an on-demand streaming media session between a client and a Real Networks streaming media server.

1. User clicks on a web link:
<http://www.startrek.com/ramgen/unimatrix-zero.rm>
2. Real server sends a .ram file that contains an RTSP link
(rtsp://www.startrek.com:554/voyager/unimatrix-zero.rm)
3. The browser forks off a Real player in a separate window (based on the MIME type received with the .ram file, e.g. x-pn-realaudio). The Real player connects with the server via RTSP (rtsp://www.startrek.com:554/voyager/unimatrix-zero.rm)
4. The server opens an RDT connection to the player and begins streaming.
5. The player sends an RTSP command to terminate the session.

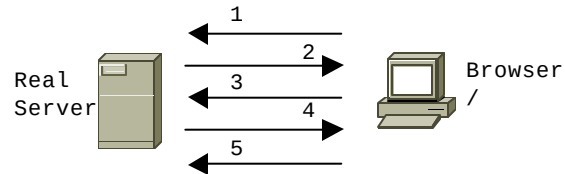


Figure 4: Client-server interactions for Real video

3. IP VIDEO DISTRIBUTION

Distribution of video over core and access network is necessary for delivery of IP video services. This distribution architecture is dependent on the type of video service offered. This is discussed below.

There are two main categories of video service: on-demand and live/scheduled. An on-demand program (such as video-on-demand or VoD) is receiver-controlled in that each receiver fully controls the streaming session. At any time, the receiver can initiate the program stream, control it via VCR-like functions (or “trick” modes) such as fast forward, rewind, pause and also terminate the program.

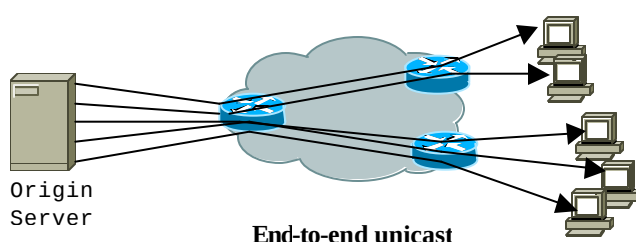
Live/scheduled programs, on the other hand, are source-controlled. A receiver can only tune to such a program (to listen to it, watch it, etc.) but has no ability to further control it. This is because live/scheduled programs are transmitted continuously by the server during pre-defined time intervals. The difference between a live versus scheduled program is primarily based on the source of content. For a live program, content source is a real-time event that is encoded and transmitted as it is happening; Content source of a scheduled program, on the

other hand, is a pre-recorded event stored on a video tape, DVD/CD, film, etc.

There are multiple ways to transmit a video program over an IP network. These include unicast, multicast, splitting (or reflected unicast) and edge multicast (or reflected multicast).

Unicast

In a unicast transmission, each receiver receives its own stream from the media server (see figure below). If there are one thousand active receivers, then there are one thousand streams being served by the media server. Though unicast allows a receiver to fully control the stream, it quickly becomes unscalable due to the significant demand it places on network and server resources. For live/scheduled programs, unicast is clearly quite inefficient.



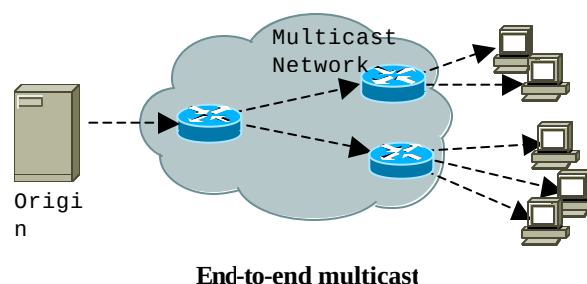
Multicast

Multicast is an efficient Layer-3 mechanism for delivering a media program to multiple receivers. Here, the server sends out a single program stream to a multicast group². The multicast-enabled IP network to which this server and the intended group of receivers are connected is responsible for delivering this stream. If the receivers are spatially distributed (in network sense), then the IP network replicates the stream as necessary to reach all receivers of that multicast group. Multicast is quite efficient for live/scheduled programs because (1) it requires minimum server resources – only one stream per

² IANA has assigned the Class D address 224.0.0.0 - 239.255.255.255 for multicast groups (i.e. destination IP addresses).

program regardless of the number of receivers and (2) minimum stream replication in the network – at most one stream per link depending on the receiver distribution. This is shown in the figure below.

Note that the server continuously transmits each live/scheduled program stream. Receivers join and leave the program asynchronously using the Internet Group Multicast Protocol (IGMPv2, RFC 2236). The IP multicast network is responsible for building an efficient tree to distribute multicast traffic. Edge network devices are responsible for maintaining the group membership of receivers and replicating streams to them.

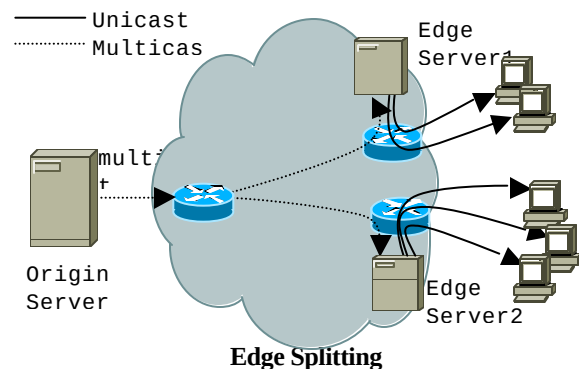
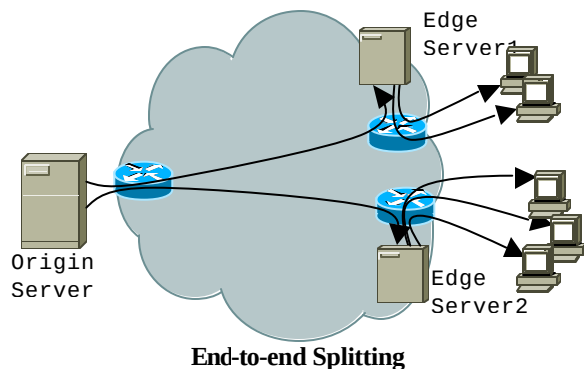


For an overview of IP multicast, visit http://www.cisco.com/warp/public/cc/pd/iosw/pr/odlit/ipimt_ov.htm. For a multicast quick-start configuration guide for routers, visit <http://www.cisco.com/warp/public/105/48.html>.

Splitting

The term “splitting” is coined by Real Networks for an efficient delivery mechanism of live/scheduled programs over a non-multicast network. In IP multicast, stream replication occurs by IP network elements whereas, in splitting, replication is performed by video servers that are placed within the network. This is illustrated in the figure below.

Splitting allows tree-based distribution of unicast streams to the network edges. Each receiver



connects to an edge server close to it (as determined by the request routing mechanism). The edge receiver then determines a path from it to the origin server. The number of program streams served by the origin server equal the number of servers participating in that program at the next lower hierarchy. In the 2nd-level hierarchy shown in the figure above, both 2nd-level (edge) servers are participating in streaming the program. Thus, the origin server is sending two unicast streams, one to each participating edge server. Each edge server splits the incoming stream into multiple unicast streams, one per each participating receiver.

Edge Multicast

Even with splitting, the edge server fan-out can be a concern when serving to a large community of users. Hence, it is beneficial to convert the edge portion of the non-multicast network to multicast. Streams traverse most of the network as unicast. The last-hop delivery of this stream, from edge to user community, is via multicast. As before, each receiver requests the edge network element to join to a multicast group via IGMPv2.

There are two mechanisms to deliver program streams over a non-multicast network: splitting and GRE tunnels. In a split-based distribution, each edge network element receives a multicast program from some edge server. Each such edge server receives the stream from the origin server through a tree-based server hierarchy.

Alternatively, the origin server continues to send a

multicast stream. GRE tunnels are created within the non-multicast network portion to encapsulate multicast traffic.

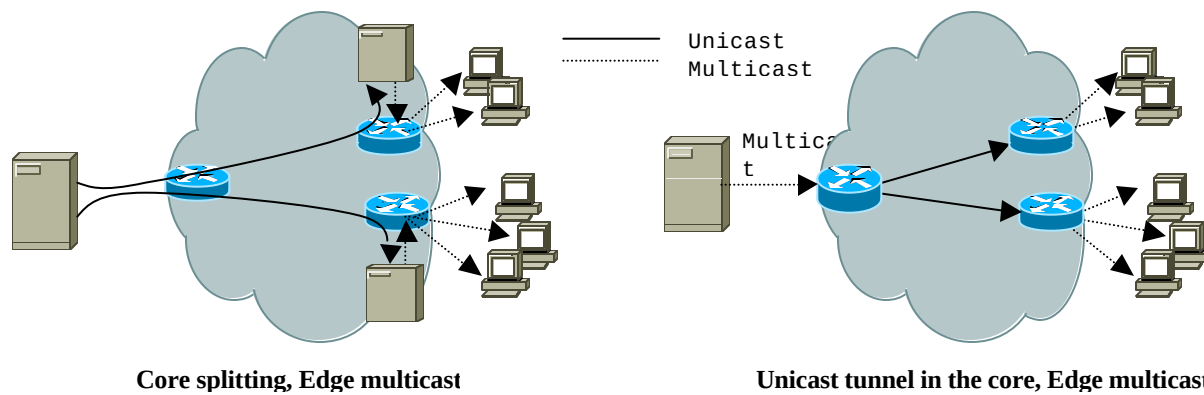
In either case, it is desirable to construct unicast delivery paths (trees or tunnels) dynamically, as per user demand. If no users are joined to a multicast program at a given network edge, then there is no need to send that stream to this edge.

4. IP VIDEO OVER CABLE NETWORK

Most two-way HFC networks today consists of broadcast video delivery via MPEG TS and HSD delivery via IP/DOCSIS. DOCSIS is assigned a separate 6MHz QAM.

Broadcast MPEG TS channels are statically assigned within a single QAM. For instance, 10 broadcast video channels, each having bitrate of 3.5Mb/s, can be multiplexed in a single QAM-256 (with aggregate bitrate of ~38Mb/s).

Peak concurrent IP video streams over DOCSIS QAM is the same as that for MPEG QAM for the same stream bitrate. One major issue with DOCSIS is that only one DOCSIS channel can be assigned to a cable modem. This limits the number of subscribers that can access video simultaneously. This is exacerbated by sharing of DOCSIS for HSD and potentially voice telephony. Dynamic channel change (DCC) will



be a critical feature for scalable deployment of data, voice and video services over DOCSIS.

Additionally, like MPEG video, IP video requires many more downstream QAMs. This implies that QAM density on cable modem termination systems (CMTs) needs to increase so that the economics of IP video delivery become reasonable.

Detailed traffic engineering analysis for multimedia services over HFC is presented in [12].

Content Distribution

In a typical cable network, the regional head-end and local hubs form a two-level content distribution hierarchy that is well suited for streaming media delivery. High-demand content is located as close to the subscriber as possible to reduce bandwidth demands of streaming media in the backbone network. Less popular content is available at the headend on an on-demand basis.

Each local hub contains a cluster of streaming media servers, caches and storage devices, collectively referred to as the edge media server. and can intelligently balance the load across these servers for optimum performance and reliability. The headend contains a core media server which typically has a several times bigger cache and storage system. In a hierarchical streaming media architecture, unicast streams are generated by edge media servers whereas multicast streams are

generated by the core media server. The headend also contains content acquisition, management and distribution systems. New content is brought in via satellite, encoded (if required) and stored in the core media server. Content management and distribution system allows a cable operator to centrally manage media content, to set policies for cached and stored content and to distribute content reliably to local hubs.

Centralized vs Distributed Architecture

In a centralized architecture, a core media server in the headend serves all cable subscribers. Two issues need to be considered in this situation: server bandwidth utilization and backbone fiber ring bandwidth utilization. To understand server bandwidth utilization, consider Table 1 below which depicts media storage and server bandwidth requirements as a function of stream rate. For instance, approximately 600 Mb/s of server bandwidth is required to serve 2000 simultaneous streams at 300kb/s.

Figure 4 shows the fiber ring bandwidth utilization as a function of data and streaming media penetration rates (same as the ones used in Tables 1-3). Both OC-48 at 2.5 Mb/s and OC-192 at 10 Mb/s rates are considered for the fiber ring. It is clear that the scenario of a centralized web cache and a centralized media cache consumes over 60% the OC-48 ring bandwidth at low penetration rates whereas this scenario is

	Stream Rate (kb/s)				
	300	500	700	1000	1500
Storage Rate (MB/Hr)	135	225	315	450	675
Storage (GB)					
100 Hrs	14	23	32	45	68
200 Hrs	27	45	63	90	135
500 Hrs	68	113	158	225	338
1000 Hrs	135	225	315	450	675
2000 Hrs	270	450	630	900	1350
Server BW (Mb/s)					
100 Streams	30	50	70	100	150
200 Streams	60	100	140	200	300
500 Streams	150	250	350	500	750
1000 Streams	300	500	700	1000	1500
2000 Streams	600	1000	1400	2000	3000
Streams/Link					
Ethernet (10-Mb/s, 40%)	13	8	5	4	2
Ethernet (100-Mb/s, 70%)	233	140	100	70	46
Ethernet (1-Gb/s, 70%)	2333	1400	1000	700	466
OC-3 (155 Mb/s, 95%)	490	294	210	147	98
OC-12 (650 Mb/s, 95%)	2058	1235	882	617	411

Table 1: Media storage, server bandwidth and number of streams served per interface link for different stream rates

acceptable for OC-192. In the case of OC-48, a fully centralized architecture works only at low penetration, a partially centralized architecture (centralized web cache and hierarchical streaming caches) works at medium penetration and a fully hierarchical architecture is required at high penetration. Of course, as stream rate increases, a fully hierarchical solution is necessary.

For hierarchical streaming media architecture, some key requirements for core and edge media servers are discussed below.

Edge Media Server: At low penetration (10% to 15% of HHP), an edge media server with interface bandwidth of OC-3 or fast ethernet is sufficient. one to two thousand hours of stream storage and ability to serve several hundred streams simultaneously is sufficient. . For higher penetration (over 40% HHP), however, the edge server must be able to scale to aggregate interface bandwidth of up to an OC-12 link (about 600 Mb/s). Stream storage capacity of 50 GB initially or about 370 hours and should

scale to to 200 GB or about 1850 hours at 300kb/s.

Core Media Server: The core server serves multiple edge servers (during cache misses) and also is the source for all multicast transmissions. . Link bandwidth of 300 Mb/s to 1 Gb/s is required, with ability to scale to multiple gigabits/sec for higher stream rates. The core media server storage capacity should be 2x to 5x more than an edge server, i.e. in the range 100 GB to 1000 GB.

Quality of Service

A critical consideration when designing streaming media networks is quality of service (QoS). To deliver a consistent, high-quality user experience, the network must maintain a given stream's specified latency and throughput requirements. Though end-to-end QoS for streaming content is generally difficult to guarantee over the public internet, it becomes more manageable over a DOCSIS-based private cable network. In the HFC network segment, from local hub to subscriber STB or cable modem, QoS is

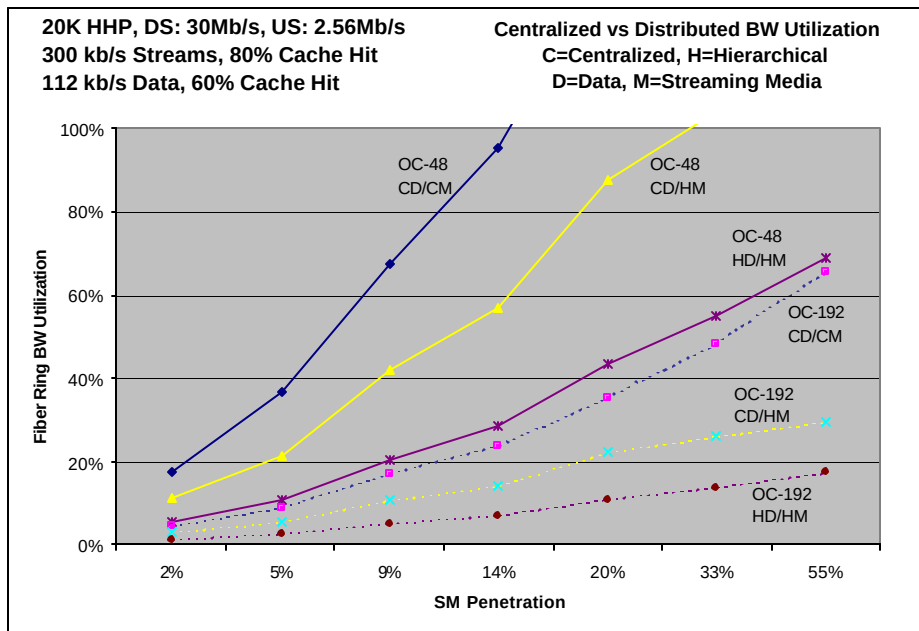


Figure 4: Fiber ring bandwidth utilization for centralized and hierarchical topologies for data and streaming media

achieved with dynamically provisioned service flows specified under DOCSIS 1.1.

On the optical backbone network segment, from the headend to the hub, QoS is achieved using either Differentiated services (Diff-serv) or Resource Reservation Protocol (RSVP) or thorough a combination of both [9,10]. Diff-serv classifies IP packets into a few aggregated classes using the type service (ToS) bits or diff-serv code points (DSCP). QoS-enabled routers and switches can then shape traffic using intelligent queuing (e.g. class-based weighted fair queuing or CBWFQ) to queue video streams as preferential class traffic ahead of other best-effort class traffic.

In addition to Diff-serv, RSVP signaling may be used to reserve bandwidth on a per-data-flow basis through the edge HFC network. RSVP signaling is used to dynamically initiate, configure and terminate DOCSIS 1.1 service flows from CMTS to cable modem. See [8] for the use of RSVP to set up voice telephony sessions under the PacketCable DQoS specification.

Set-top Boxes

A key challenge in delivering IP video to TVs is the availability of a suitable STB. As indicated earlier, the three streaming media formats used today are of proprietary nature. If decoding were to be done in software, significant CPU resources and run-time memory may be needed. Also, integration of the streaming video player with the STB middleware GUI is required.

Some next-generation STBs may contain programmable video processors (e.g. Philips TriMedia processor) that support multiple codecs in firmware. This will allow the support of MPEG-x and popular proprietary codecs in an economical manner.

Conditional access system (CAS) and digital rights management (DRM) are also necessary to protect copyrighted content and to provide differentiated video services. An open CAS/DRM mechanism for IP video is required for a cost-effective solution.

REFERENCES

- [1] Broadband Net access nearly 8 million strong, Corey Grice, CNET News.com February 28, 2001, <http://news.cnet.com/news/0-1004-200-4983492.html>
- [2] H. Benoit, *Digital Television*, John Wiley, 1997.
- [3] Michael Adams, *OpenCable Architecture*, Cisco Press, 2000.
- [4] RFC 1889, RTP: A Transport Protocol for Real-Time Applications, Internet Engineering Task Force (IETF), January 1996.
- [5] RFC 2502, Compressing IP/UDP/RTP Headers for Low-Speed Serial Links, Internet Engineering Task Force (IETF), February 1999.
- [6] RFC 2250, RTP Format for MPEG1/MPEG2 Video, Internet Engineering Task Force (IETF), January 1998.
- [7] RFC 2326, Real-Time Streaming Protocol (RTSP), Internet Engineering Task Force (IETF), April 1998.
- [8] PacketCable Dynamic Quality of Service Specification, <http://www.packetcable.com>
- [9] RFC 2205, Resource ReSerVation Protocol (RSVP) – Version 1 Functional Specification, Internet Engineering Task Force (IETF), September 1997.
- [10] RFC 2475, Diff-serv, Internet Engineering Task Force (IETF).
- [11] Ivy Lui and Prashant Gandhi, Streaming Media Opportunity for Cable, Cisco white paper, December 1999, http://www.cisco.com/warp/public/cc/so/cuso/sp/stmda_wp.pdf
- [12] John Chapman, Multimedia Traffic Engineering for HFC Networks-White Paper, Cisco Systems, 1999. http://www.cisco.com/warp/public/cc/so/cuso/sp/hfcn_wp.pdf

A SECURE DELIVERY AND USAGE MODEL FOR RICH CONTENT THAT MINIMIZES PEAK NETWORK CONGESTION

Walter Boyles
Wavexpress, Inc.

Abstract

Broadcast is the most efficient means of delivering rich content to the consumer. Unicast pull is the most convenient method for the consumer to obtain rich content on-demand (any kind of content at any time). The inherent shared bandwidth nature of the Hybrid Fiber/Coax (HFC) network, coupled with advances in technology, allows the best attributes of both broadcast and unicast pull to be provided to the consumer with the lowest cost and a high degree of consumer choice and discretion. With this, reliable security and transactional methods must be used to protect the content and verify consumption.

Overview:

Broadcast either to the set-top-box (STB) or the cable modem can be utilized in addition to and, at times, as an alternative to on-demand delivery. For set-top box delivery, video-on-demand (VOD) allows the consumer to select, from a menu of movies, a title thus dedicating bandwidth to that consumer for the duration of the film. However the terms of sale, typically, will also allow the consumer to pause that film for up to a specified period of time (i.e. the film must be viewed within a 24-hour time period). The MSO can choose to maintain that dedicated bandwidth during the pause or reclaim it for use by other viewers. However, even if that bandwidth is reclaimed, the terms of sale frequently will mean that the MSO must be able to provide alternative bandwidth for a return to viewing in a reasonable period of time after the pause by

the consumer. As such, reserve bandwidth must be allocated to make certain that the consumer does not return from a pause and the terms of sale cannot be met by the MSO.

VOD is not yet widely deployed, but it can be anticipated that peak periods for VOD usage will exist (e.g. Saturday evenings from 7 PM – 12 AM) and that slack periods will also occur (e.g. Monday mornings from 7 AM – 11 AM). Bandwidth utilization for VOD will be extremely uneven.

For cable modem traffic, peak periods of utilization may similarly occur in evenings when consumers will experience noticeable degradations in performance due to bandwidth demand within the HFC network and/or to delays outside of the network (e.g. many hits on a server) as compared to early mornings when demand within the network as well as overall Internet traffic may be less. In either case, a way to equalize loading on the network over time provides benefits to consumers and to operators alike.

Storage Technology:

Over the last ten years, while Moore's law has become a well-known and widely quoted phenomenon, the time duration between the doubling of chip densities (the number of transistors on a silicon substrate), has been slowly increasing to upwards of 18 months.

Meanwhile, advances in magnetics technology have allowed for the doubling of hard drive densities in half that time (approximately nine months). During the same period, while the cost of hard drive

storage has been decreasing rapidly, the last mile bandwidth available to the average consumer has been increasing relatively slowly. Even for those fortunate consumers with broadband connectivity, the price and the speed of their connection has changed little over the past few years.

Delivering content in broadcast or multicast modes, that is subsequently cached, can allow consumers faster access to that content. This type of delivery is, in part, enabled by large amounts of inexpensive storage being made available to the consumer.

Year	1988	1998	1999	2002 (predicted)
Cost/ MB	\$11.54	\$0.04	\$0.02	\$0.003

Table 1. Hard Drive Density vs. Cost (Disk/Trend).

Table 1. shows industry mean storage prices. However, consumers buying desktop PCs typically buy in the most economical segment of the hard drive industry. Today for mid-range desktop systems, the cost of storage is often already at \$0.003 per incremental MB.

Additionally, with drive densities doubling every 9 months, the cost of storage in the lowest cost part of the curve, in early 2004, may be extrapolated to be as low as one hundredth of a cent per MB. That calculates to a cost of approximately 35 cents to store 3.5 GB video:

$$3500 \text{ MB} \times \$0.0001 / \text{MB} = \$ 0.35$$

Without massive increases in last mile bandwidth to the consumer, the trend will be that it is cheaper to store large files than to transmit them in unicast.

And with drives approaching 200 GB already available and mid-range desktop systems typically equipped with 20 - 80 GB drives, storage already greatly exceeds the

space taken up by the operating system and typical applications software packages. Further, again extrapolating drive densities forward we can expect that by the end of 2002, drives exceeding 1,000 GB (1 Terabyte) will be available and new mid-range consumer systems should nominally have hard drives of 100 - 400 GB. This reduction in the price of storage will equally allow the addition of large amounts of storage to be added to set-top boxes.

Now, as with Moore's law and the increase of densities of integrated circuits, it is well recognized that increases in drive density eventually run into barriers imposed by the laws of physics. In the case of magnetics, one barrier is the supermagnetic effect (SPE). The SPE limitation will occur when the amount of energy of the magnetic spin in the atoms, that constitutes one bit of information, approaches the ambient thermal energy. When this occurs, the bits become subject to random "flipping". It is believed that the limitations imposed by SPE may limit miniaturization as early as 2005 at around 150 Gb per square inch.

Beyond this 150 Gb per square inch limit, new strategies are already being devised. These include techniques from changing the orientation of bits to new magnetic materials and even to moving towards the addition of optical materials to magnetic materials. With these kinds of techniques it is entirely possible to conceive of consumers possessing many terabytes of storage on their PCs in the not so distant future.

With the amounts of storage becoming available in the shorter term, the delivery of major interest movies to all subscribers on a network (who elect receipt) will be possible while allowing the same viewing flexibility as VOD with a single uninterrupted transmission by the MSO. Similarly, one could allocate a portion of the 38 Mb/s bandwidth (based on

256 QAM) of a 6 MHz cable modem channel to broadcast during off-peak hours. Delivering 5Mb/s during the 2 AM- 7 AM period of low-usage would mean the availability of over 100 GB of content to the cable modem customer:

$$(5 \text{ Mb/s} \times 3600\text{s/hr} \times 5 \text{ hr}) / 8 = 112.5\text{GB}$$

The consumer would have the ability to elect to cache any portion or all of that day's delivery of content.

This would allow the consumer to spend a portion of his viewing time (most likely in the peak usage period of the evening) consuming that cached content and thus further limiting peak usage. Thus, the use of caching content received from off-peak delivery can result in the usage of excess bandwidth in slack hours and decreased demand during peak hours. The bandwidth available and bandwidth demand are then more equalized and the network operates closer to maximum efficiency.

Security:

In delivery of stored content, a number of security mechanisms become extremely important. First, a reliable system of copy protection for content cached to the drive of the PC or the STB is needed. Second, a security system that allows protection of royalties for content consumed is required. Third, a reliable and cost-effective method of transacting on the content purchased, leased or viewed is a key element.

To satisfy these requirements, the security of the content, the permission to access that content, and even the transaction on the content itself are most reliably achieved through a hardware security device. This conclusion is consistent with the methods used today by the cable industry in allowing access to premium content where security is contained within the STB. Monolithic solutions are, of course, preferred that do not allow hackers access to any bus containing cleartext data.

The rationale for hardware security involves many factors but to a degree it comes down largely to one. Software is relatively easy to tamper with and software hacks are easy to replicate on a large scale. The delivery and consumption of content is far different from an Internet purchase of a physical good where security involves encrypting a credit card number for a purchase that the consumer is authorizing. In content security, there is a temptation and an opportunity to avoid paying for something in an area where precedents like unprotected MP3s already exist.

A hardware device that allows some degree of programmability and renewability is also preferred. First, a programmable device can provide multiple security functions including persistent protection of content which is vitally important in an open system like a PC. Second, a system with some degree of renewability can utilize a smart card or preferably, for true hardware renewability, a more complex portable token that avoids some of the limitations (e.g. interface and performance) of an ISO 7816 smart card. If smart cards are used, care must be taken to avoid unsecured smart card readers in a PC environment since key security inside the PC is essential to maintaining content security.

The requirements for security in a PC for copy protection are different than in consumer electronic devices. In consumer electronic devices, a bus-oriented copy protection system, between say a DVD player and a monitor, may appear adequate. However, in a PC, the variety of ways to access that content inside the PC plus the number of ports make content protection a more complex issue than in the consumer electronics space.

Digital rights management systems (DRMs) have proliferated in the PC space although their actual usage has been somewhat limited due, in large part, to consumer behavior and lack of consumer

acceptance. However, software DRMs still exhibit the same security weaknesses as other software solutions do and provide less resistance to tampering than well thought out hardware solutions including hardened DRM.

Additionally, since content delivered and cached will typically be viewed on a time-shifted basis, the ability to securely transact this content will be also be of importance. Ideally, a set of security mechanisms and methods for processing transactions off-line would provide the benefit of not requiring repetitive server transactions. Content consumed and paid for, in this way, would provide the benefit of minimizing operator cost in the same way as the single delivery of a copy of content to multiple consumers in the network did. The decryption and transaction are most reliably secured in a single tamperproof hardware device attached to the PC or embedded in the set-top box.

Since the cable modem-attached computer receives content over multicast and frequently the interactive STB has a back channel available, a strong key exchange method is required. The security mechanisms (including copy protection) should avoid use of hidden global secrets or other mechanisms, which may result in catastrophic failures.

Users and preferably legitimate devices (both) should be authenticated by a reliable process that disallows common known attacks. Authenticated users should be able to prove that they possess legitimate devices for metering and decryption of content.

Unique IDs are widely used for to identify devices such as the IEEE 802.3 48 bit addresses in NIC cards. However, for maximum security, IDs should not only be immutable (not able to be reprogrammed by the hacker) but shielded (secret) from the consumer. Additionally, legitimate hardware IDs should not be able to be guessed by potential hackers.

One mechanism for identifying legitimate devices is the assignment of non-deterministically, random identities at the time of manufacture of devices. By assigning immutable secret identities, of adequate length to devices, illicit clones without legitimately assigned identities are detected. The attacker's ability to guess one legitimately assigned random number represents a problem comparable to exhaustive key search for a strong symmetric key cryptographic algorithm. One issue is the potential for the occurrence of multiple instances of any one legitimately assigned random ID.

For example, if a 128 bit random identity is assigned, this constitutes a total of 3.4×10^{38} different random numbers. However, if two hundred million consumer devices (2×10^8) devices are built the probability that no two of these devices possess the same random number is:

$$n = x - 1$$

$$\prod_{n=1} \left(1 - \frac{1}{y} \right).$$

Where x = the number of devices and y = the number of possible random numbers. Therefore, the probability that there are multiple instances of one ID is:

$$n = 199,999,999$$

$$p = \left(1 - \prod_{n=1} \left(1 - \frac{n}{3.4 \times 10^{38}} \right) \right) > 0.$$

Thus, there exists a possibility of multiple instances of the same ID and one cannot automatically assume the existence of multiple instances constitutes an illegal clone.

The binding of device identities to user identities (the user identities may be kept in a smart card or other portable token) can be

used as a barrier to useful theft of the device. However, the value of having a user identity bound with a random device identity at the time of purchase also provides a more effective barrier to random guesses. That is, the attacker (guesser) must not only guess an assigned (secret) device identity but also know the user identity to which it is bound. Many guesses with the same user identity are easily detected. Also, since there is no longer sole reliance on device IDs, the problem of multiple instances of the same legitimately assigned device identity can be minimized.

Once a legitimate device and user are determined, a user's security device can participate in a public key exchange and the sending of the symmetric keys to allow the access to content. In some cases, the keys transmitted will not decrypt the content itself but rather other keys that in turn are used to decrypt the content. Changing keys is done on a periodic basis; in some cases, systems require that many key changes be made in a relatively short period of time.

The decryption of the cached content should occur subsequently or coincide with a payment for that content. Payments can occur either on-line or off-line. Secure off-line transactions have the advantage of minimizing server usage and often also transaction fees. Off-line transactions, of this nature, are typically debits and additional security, for the transaction being recorded, can be provided by a real-time clock in the security/debit mechanism.

Once content is transacted on, a method of usage must be employed to prevent unauthorized redistribution of content. In a closed system, a local usage method may be employed or instead one where a smart card accompanies the content with a consumer to other devices/players. In an open system such as a PC, persistent protection by some system such as digital rights management technology may be employed.

The use of a programmable security processor can allow purchase, decryption, and persistent protection inside the same device. Devices of this nature must have a well-conceived trust assurance network for loading code. A secure operating system within these devices may also be desirable.

A programmable security processor will contain secure volatile and nonvolatile memory. The limitation of a monolithic solution containing volatile and nonvolatile (i.e. for key storage) memory is often the amount of volatile memory (typically SRAM) that can be placed economically on-chip. This limitation is due to the processing difficulties of fabricating custom ICs with both DRAM and nonvolatile memory (e.g. Flash or E²PROM).

In the future, devices of this nature may be fabricated with ferroelectric memory (FRAM) that may take the place of both the volatile and nonvolatile memory blocks thus allowing massive code spaces on-chip. This type of chip architecture would allow very complex content usage and rights management software to be loaded entirely inside a single chip thus maximizing security.

Conclusion:

The cable industry has continued to innovate products and services to the consumer. For consumers, the movement towards cable modems and digital cable, is enabling access to much more content with a greater degree of ease than ever before.

Movement to more efficient modulation and compression techniques coupled with advances in architectures allows MSOs to more effectively use the capacity of their system. However, the addition of delivering content during off-peak hours in broadcast to the STB or IP multicast to the cable modem

consumer's hard drive enable even greater efficiencies by the MSO and more choice and greater access to the consumer.

The underlying technologies, that allow business models which employ off-peak delivery and caching of encrypted content for subsequent consumption and off-line payment, are increases in storage, strong security mechanisms including copy protection, and reliable methods of transacting on the content that is to be consumed.

Bibliography:

Data-Over-Cable Service Interface Specification, Cable Modem to Customer Premise Equipment Interface Specification, Cable Television Labs, 2000.

Dell, <http://www.dellforme.com>, 2001.

Huitema C., IPv6, Prentice Hall, 1997.

Toigo, Jon W., Avoiding a Data Crunch, Scientific American, May 2000.

ABUSE AND FRAUD DETECTION IN HIGH-SPEED-DATA NETWORKS

Why and How can HSD Operators deal with and profit from abusers

Pat Darisme, NARUS Inc.

Abstract: As MSOs¹ are increasing penetration of 'always on' high-speed-data (HSD) service, some zealous users are abusively taking advantage of this low-cost, high-powered flat rate service. Some MSOs are painstakingly finding users who are running small ISPs from behind their residential cable modem, using a NAT box to share a single IP between multiple systems, running residential or commercial web servers, spamming mail servers or newsgroups or performing malicious attacks on other users. These users are expensive both in terms of the network resources (network bandwidth, CPU cycles, hardware upgrades) and with the man hours consumed in identifying & dealing with these users. Most of these issues are likely to remain problems even with the advent of DOCSIS 1.1. This paper will explore various types of abuse and methods of dealing with them.

INTRODUCTION

The Internet started as a collection of academic and research networks, but since 1994 it has become a key part of North-American life in over 149 million households^A. The advent of high-speed residential Internet access via Cable and DSL is providing residential users with 'always on' high-speed connections to the Internet. For a large percentage of subscribers these connections bring the known benefits of broadband (more bandwidth, lower latency, always on, no busy signals etc.). These connections also

open the door to new types of abuse from a small percentage of customers.

Abuse in the context of HSD can be loosely defined as any behavior by a subscriber which violates legal statutes, or service provider policies, including activities which cause the subscriber to not pay for all the services consumed or resources (network, server, CPU etc.) consumed or to impact the quality of service of other network users.

As the Cable Industry sees an increase in penetration of HSD service (some large systems have passed the 20% penetration mark) it finds itself offering Internet access to a large number of North-American Internet users: over 5 Million in 2001 and over 19.5 million expected by 2005^B. Along with a growing piece of the revenue stream comes increased necessity to police abusers.

This paper will address three questions:

1. Why identify and deal with Abusers
2. Who are the Abusers
3. What type of Abuse exists and what solutions can be applied

1) WHY IDENTIFY AND DEAL WITH ABUSERS?

HSD service providers need to identify abusers not simply to be good corporate citizens or cybercitizens but also because abuse increases operating costs, can decrease revenue and degrade the quality of service (QOS). The diminished QOS can lead to bad press and damage to the service provider and MSO's brand.

¹ MSO – Multiple System Operator a.k.a. Cable Operator

A. Operating Costs

Operating a high-speed data network involves high capital costs (routers, servers etc.) as well as high operating costs (bandwidth, staff etc.). When a few abusers increase network usage this can lead to increased capital costs in the form of network capacity upgrades, or higher bandwidth costs, both of which lengthen the path to profitability.

In a recent example, a service provider who offers a free hosting service as part of their product offering was concerned about growing bandwidth consumption and costs. They purchased an abuse detection system relying on real-time traffic Analyzers and an Internet Business Infrastructure Platform. Using this new system, customers who were abusing the free service by distributing large files such as MP3 and 'Warez'² files were identified and removed. The peak bandwidth consumption of their hosting service dropped from 600Mbps to 20Mbps leading to savings estimated at \$150k/month (Their bandwidth cost was approximately \$250/Mb)

B. Quality of Service (QoS)

Network abuse reduces the quality of the Internet service customers are paying for. Abuse consumes resources such as limited HFC (Hybrid Fiber Coax or Cable Plant) bandwidth, costly peering bandwidth, router and server CPU cycles, storage etc. and can even cause partial or full network outages. Limiting this abuse can reduce the frequency of operational problems, limit the necessity of system upgrades and improve the overall quality of the product offered.

² MP3 files are audio files (typically music) which use the MPEG 3 compression standard. Warez is Internet jargon for pirated Software

o Quality of Service & DOCSIS 1.1

Cable Operators sometimes mistakenly think that DOCSIS 1.1, with its ability to implement rate limiting, will resolve all QoS issues. With DOCSIS 1.1 on the HFC, service providers still need to over-subscribe the bandwidth available both on the HFC and on the back end IP network and peering links. This is done to save costs because service providers expect subscribers will not all be consuming their allocated bandwidth at the same time, or all the time, and because of the burstable nature of IP traffic. While the subscriber may be entitled to only a rate-limited 128 kbps, data models do not expect that the subscriber will use this 128 kbps 24 hrs/day 7 days/week – yet this is what some abusive subscribers do by leaving streaming audio/video services on, or running Napster or other servers, on their 'always on' connection even in their absence. With the deployment of DOCSIS 1.1, HSD providers, will still be subject to various forms of abuse.

C. Bad Press

Abuse can cause a network outage, network delays (email delays are particularly popular), or security problems all of which can lead to bad press and dilute the MSO's brand. The Press seems to enjoy articles on network outages, especially when these occur at large brand name carriers. These negative articles affect the public's perception of the quality and value of the HSD service.

D. Revenue loss

The 2-way nature of high-speed-data makes traditional theft of service practically impossible. However, in a flat rate billing environment, it is possible for paying subscribers to use more services than they are paying for. This results in a loss of potential new revenue for the operator.

E. Peering

Service providers (or domains) who harbor abusers or who do not take steps to limit abuse on the Internet, risk having their peering connections cancelled by upstream service providers, having services shut down (e.g. Usenet Death Penalty) or having their domain added to Blackhole lists which filter traffic. This eliminates their ability to offer Internet access³ or services (email / news). Limiting abuse can allow the service provider to maintain the network peering relationships and service feeds (such as email & NewsGroups) which are key to offering HSD services.

F. Law Enforcement Compliance

Some jurisdictions have regulations which may require service providers to have the capability to monitor the Internet activity of selected subscribers when requested. Having proper systems in place to track these subscribers may reduce the operational burden of compliance with these law-enforcement mandates which can come on very short notice and offer very little latitude.

G. Open Access / Multiple ISP Networks

Abusers can cause problems in traditional exclusive access HSD networks where the MSO bears both the full revenue and the burden of abusers. In this environment the MSO has full control over the customer and can either charge abusers for the usage, or

³ Though the canceling of peering links because of abusive activities (attacks, spam etc.) in a domain is somewhat rare this is mostly due to the fact that dedicated network operations' staff have spent many long hours combing through network logs to eliminate abuse before being cut-off. This reactive approach drains resources which are better suited to designing and scaling networks.

shut them off. In Open Access or Multiple ISP HSD deployments the MSO may be getting only a small amount of revenue from abusers who are consuming disproportionate amounts of shared HFC bandwidth or services. In this Open Access environment, the MSO may not control the customer and the ISP may not be interested in co-operating with the MSO to identify abusers. Even if contracts ensure the ISP will deal promptly with abusers, compliance is more likely if the MSO has the ability to verify abuse. In the end the burden of protecting the limited resources of a shared HFC infrastructure rests solely on the shoulders of those who have invested in it and who must protect the quality and investment already made in their brand, the MSO.

There are multiple business reasons why a service provider needs to identify network abusers. Despite the issues described above, service providers often either ignore abuse issues or have a small team of dedicated staff who struggle in a reactive mode to contain damage rather than use the right tools to proactively prevent abuse.

2) WHO ARE THE ABUSERS?

There are different types of subscribers who can abuse the network. They can be loosely categorized as follows:

A. Hackers

Hacker is a term used by some to mean "a clever programmer" and by others, (especially journalists or their editors), to mean "someone who tries to break into computer systems." (The Internet community tends to call those who break into computer systems 'Cracker'.) Hackers are typically power-users who have a very good understanding of computer systems, protocols and programming, and who find great joy in learning more about and using the

power of computers. Some would argue that the Internet, Usenet and Unix/Linux were created and are maintained by hackers, whose skills can make them desirable engineers. One of the best-known hackers/crackers is the skilled and infamous Kevin Mitnick who was jailed for over 5 years following some computer attacks in 1994.

B. Script Kiddies

These are novice users who would like to be known as true hackers but who resort to downloading or copying simple programs posted on the Internet. 'Script Kiddies' run these scripts - often without understanding the exploit - either out of curiosity, in hopes of achieving notoriety, or to get even with someone they are unhappy with. The downloaded programs, or scripts, are automated ways of generating large amounts of traffic or other network disturbances; they affect the victims' network connectivity and potentially result in an outage. One recently well known 'Script Kiddie' is the 15 year old - identified as 'MafiaBoy' - arrested for the attacks on popular internet sites such as Yahoo, eBay and Cnn.com in the spring of 2000.

C. Unwilling victims (Trojans)

Subscribers can unknowingly generate unacceptable traffic when their computer becomes compromised either by downloading software which contains a virus or exploit, or when an intruder places such software on their computer following a break-in. This software can be set to generate or relay large amounts of data following some trigger (time, program execution, external trigger etc.).

D. Fraudulent Businesses

A subscriber can be running a fraudulent business such as a pyramid scheme or some other 'make money fast' system, such as selling cable decoder boxes or providing illegal services. They may have subscribed to high-speed data services for the sole purpose of running their business from home.

E. Criminals and Terrorists

HSD subscribers are not unlike most other segments of the population, which means that a certain percentage of subscribers are also engaging in criminal activities and using their Internet connection as a means of facilitating or perpetuating these activities.

F. Regular Subscribers

Some level of 'abuse' will simply be due to regular subscribers. These may be power-users who are ideal candidates for a premium service offering; or telecommuters who have use that is much higher than average or that violates AUP guidelines (perhaps by connecting multiple PCs to a single Cable Modem connection).

3) WHAT TYPE OF ABUSE EXISTS AND WHAT SOLUTIONS CAN BE APPLIED?

In the following sections we will explore the different types of network abuse, and for each, the type of people who cause it, its impact and different service provider solutions.

A. Spam or excessive messaging

'Spam' is Internet jargon for what is also known as Unsolicited Commercial Email

(UCE). This is loosely defined as multiple⁴ messages with substantially the same content. These messages can either be email or newsgroup postings. They can be destined to one or more recipients or newsgroups⁵. Typically they have subject lines such as 'Make Money Fast' or 'Free Cable TV' or 'Free Porn'. Recently there has been an increase in email viruses⁶ which spread quickly through the Internet, congesting mail servers at service providers and corporations. These can also be considered Spam.

- Who – 'Fraudulent businesses' using Spam for marketing. Hackers / Script Kiddies or even regular Subscribers may also - out of malice, vengeance or desire for recognition - generate excessive messaging.

Messaging problems can be caused by a subscriber within the Service Provider's network or can originate on the general Internet.

- Impact – Quality of Service (QoS), Bad Press & Cost, traffic filtering
These multiple messages congest servers (mail & news), consume precious CPU cycles, fill storage space, and consume network bandwidth. Messages are queued as the limited servers struggle to process their increasing number. Delayed or lost messages that ensue, lead subscribers to become quite vocal. On some occasions the press may even be alerted to the problem, causing potentially negative exposure to the brand and the underlying product.

⁴ Though the exact amount of "multiple" or 'n' hasn't been clearly defined, partly because violators would then simply send 'n-1' unsolicited messages, an generally accepted number is 20.

⁵ Cross-postings or a same message posted to multiple newsgroups are typically called 'ECP' or Excessive Cross Postings and are a form of Spam.

⁶ 'I Love You' and 'Naked Wife' are the most recent email virus subject lines

These problems generally lead the service provider to scale the messaging system by adding hardware and bandwidth to meet the ever-increasing load. This is a costly way of dealing with abusers in the HSD business where capital costs of servers are significant and flat rate pricing provides for unlimited usage.

A Service Provider who is often a source of Spam may see their domain (it's associated IP addresses) added to popular blackhole lists⁷. This means their email, and perhaps other Internet traffic, gets filtered and will not reach recipients in multiple domains. If an HSD provider's domain becomes known as the source of Spam it can rapidly be added to these blackhole lists. This happened to @Home's domain on various occasions according to Cathy Wittbrodt who was Director of Routing Engineering at @Home Network from 1996 to 2000.

Unfortunately the damage is already done when subscribers complain, servers are filled to capacity, the domain is 'blackholed' or when the press calls asking for comments. A reactive approach can only attempt to contain or repair damages. For these reasons and more, it is advantageous to proactively prevent Spam.

Service Provider Solutions

- To defend from Spam originating on the general Internet, Service Providers can purchase a third party mail service⁸ which filters out known Spam domains. These third party services do help to limit Spam from the Internet but do not address Spam generated by

⁷ MAPS - Mail Abuse Prevention System LLC (MAPS) [www.mail-abuse.org] is the maintainer of one such mail list that approximately 40% of all Internet Addresses use to filter out known Spam domains.

⁸ BrightMail [www.brightmail.com] provides one such service.

internal subscribers, which requires prompt internal response to prevent it from congesting the network.

If the ‘Spam’ originates on the HSD provider’s network there is more flexibility and need to deal with the problem. Once the abuser is identified, if the source is on the Service Provider’s network, the subscriber can either be eliminated or moved to a more appropriate billing plan.

Subscriber Group	Username	Traffic (MBytes)
RoadRunner	40:de:59:02:9a:b5	3,135,394.37
Juno	c0:12:29:68:90:32	1,231,152.56
WorldNet	0b:e3:11:94:8e:f0	832,742.83
Excite@Home	90:c0:77:8b:34:bc	432,479.40
Earthlink	ba:e0:29:68:8fe2	203,118.88
AOL	23:12:a0:f9:ca:00	129,418.00
Totals		5,964,306.04

CSV Export

The report results list the top 500 newsgroup users by newsgroup traffic. A newsgroup user is defined as a user with newsgroup (NNTP) traffic for the selected interval.

Figure 1: Abuse Detection Report showing users generating over 3 GigaBytes of NewsGroup traffic over a 7 day period.

Real-Time Analysis

Communications software can monitor the number or recipients of messages (email/news) sent by subscribers, while ignoring the content in order to maintain privacy. Real-time software can generate an alert when a threshold is met allowing the service provider to act immediately once the n^{th} email is detected, rather than waiting for an outage to be created when excessive damage is done.

Mediation software

Mediation can also be used to threshold the number of messages, such that messages beyond a threshold can be billed (e.g. 50 free emails/week, \$1 each additional email) effectively turning the abuse into revenue. This approach acts as a deterrent to excessive messaging, while potentially increasing revenue.

B. Reselling of bandwidth

Bandwidth resale involves the use of a single subscriber connection (single IP address, flat fee service ~\$44.95) for more than one PC. Examples of this are offering ISP services to neighbors or sharing the connection between multiple PCs, typically using a NAT server⁹.

○ Who – ‘Hackers’, ‘Script Kiddies’, Home based businesses, college students and corporate telecommuters are the types of subscribers who will purchase an unlimited bandwidth single connection and share it amongst multiple PCs. With the decrease in PC prices and the availability of home networking equipment, typical home users are also beginning to connect multiple PCs behind a single cable modem.

○ Impact – Quality of Service, Cost, Revenue

When a residential connection is shared by multiple systems it typically leads to higher data-consumption than is anticipated in the traffic and data modeling assumptions that go into capacity planning of the network. The Quality of Service drops as limited bandwidth (either HFC or peering) is consumed by this higher than expected usage. Costs can also increase as additional capacity is purchased either in the form of additional spectrum allocated to the HSD service, or additional bandwidth purchased from the upstream ISP. Since premium service offerings (offering multiple IP addresses) are not purchased, revenue also suffers.

⁹ NAT [Network Address Translation] servers use a single IP address (provided by the Service Provider) to allow an unlimited number (typically up to 253) of additional IP Clients with individual IP addresses to share that single IP address and it’s bandwidth. NAT servers can be a Linux or Windows server running special software, or be included in firewall / router systems which connect behind the Cable Modem. NAT servers are sometimes called ‘proxy’ servers.

○ Service Provider Solutions

Since NAT servers mask the multiple IP addresses of their clients they are very difficult to detect. NAT servers leave no traces on servers rendering log file parsing pointless. There are only two options for service providers to deal with NAT servers:

Usage based billing - Sharing a network connection is only profitable to customers because of flat rate billing models. Usage based billing models immediately defeat this type of abuse, or more properly align network expense to revenue generated. Much like people would never consider sharing their Cell phone or long-distance service which are billed on usage, an HSD bill based on the bandwidth or services consumed is less likely to be shared.

Real-Time Analysis – Analyzers monitoring the patterns of IP streams can potentially identify certain signatures of IP streams which are carrying IP content destined for different PCs'. Prototype NAT detection systems relying on these signatures have been designed. Though expectations are positive, these have yet to be tested.

C. Residential Server Hosting

Subscribers can use a residential HSD connection to run servers of any kind. The most frequent type of server encountered is an HTTP server, however many other servers including FTP, Mail, Telnet, news, streaming video, Peer-to-Peer (Napster) and others can generate upstream bandwidth. Home DHCP servers can serve the wrong IP addresses to other subscriber PCs causing isolated outages.

○ Impact – Quality of Service, Bandwidth, Revenue, Outage

QoS & Bandwidth - HSD systems are deployed with the assumption that bandwidth consumption is asymmetrical with much heavier downstream usage, as such downstream capacity on the HFC portion of these HSD systems far exceeds upstream capacity. Servers violate this assumption congesting the upstream bandwidth leading to degradation in the quality of service. Even with QoS policies like DOCSIS 1.1's rate-limiting, home based servers can still generate large consistent usage (albeit with reduced throughput) as the servers are accessed at all hours from the always active Internet population, rather than following the usage patterns of a more 'typical' subscriber.

According to an @Home abuse manager, an MSO used active polling to check every connection in a 20,000 subscriber system for the presence of home servers. "Out of 20,000 subscribers they found 969 servers consuming 34% of the system's bandwidth".

Revenue - Some HSD providers offer a more expensive or premium product for server hosting; unfortunately subscribers are unlikely to sign up for the pricier version when the entry-level product provides the same functionality. Even with DOCSIS 1.1 rate limiting, many subscribers are likely to settle for the expected 128k rate-limit for their home-based servers.

Outages - Home DHCP servers cause isolated outages for neighboring subscribers who may receive the wrong IP address and thus not be allowed on the network.

○ Service Provider Solutions

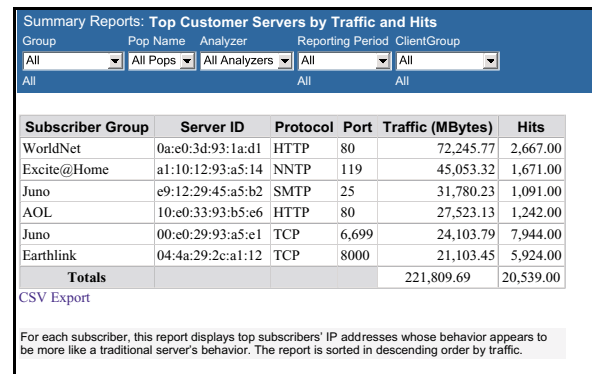
As home based servers do not leave any log file trails on service provider servers they can be difficult to detect but there are 3 options to deal with them.

SNMP – standard SNMP MIB (DOCSIS & MIB II) data collected from the modem can give a count of the number of packets sent and received by the modem (since the last modem reboot¹⁰). Although a high amount of sent packets (some call this a ‘top talkers report’) may be a hint of a home server, it could also simply be someone sending large email messages (perhaps containing pictures of their new-born) or uploading files to their hosted service someplace on the Internet (i.e. Homestead or GeoCities community web pages). Also, a ‘top talkers report’ gives no indication as to the port or service on which this potential home based server is running. SNMP data collection also places additional bandwidth on the network, potentially increasing congestion.

Active Polling – A system can be dedicated to actively poll successive ports from every subscriber IP address on the network in search for a server. Given that a server can run on any TCP port between 1 and 65,535¹¹, many successive polls (over 65 thousand polls per subscriber) are required on a typical system to identify servers. Using active polling generates traffic on the network and may well return a false negative if the polling packet receives no response for various reasons which can include a dropped packet, a server which is periodically shut-down or a server which is configured to be unresponsive to certain IP addresses or request formats.

Real-Time Analyzer – A real-time Analyzer that can see all IP sessions can easily identify the top home servers both by bandwidth and by hits, while also

identifying the port on which these are running. This gives the Service Provider a clear list of the top residential servers, as well as the necessary information to go after them and either up-sell them to a higher service, or insist the subscribers shut down the servers.



Subscriber Group	Server ID	Protocol	Port	Traffic (MBytes)	Hits
WorldNet	0a:e0:3d:93:1a:d1	HTTP	80	72,245.77	2,667.00
Excite@Home	a1:10:12:93:a5:14	NNTP	119	45,053.32	1,671.00
Juno	e9:12:29:45:a5:b2	SMTP	25	31,780.23	1,091.00
AOL	10:e0:33:93:b5:e6	HTTP	80	27,523.13	1,242.00
Juno	00:e0:29:93:a5:e1	TCP	6,699	24,103.79	7,944.00
Earthlink	04:4a:29:2c:a1:12	TCP	8000	21,103.45	5,924.00
Totals				221,809.69	20,539.00

Figure 2: Sample Customer Server Detection Report listing top customer servers with related ISP (Open Access environment) CM MAC address, port and traffic.

Usage based billing – Home-based servers increase the amount of bandwidth (mostly upstream) used by a residential connection. A Billing model based on the bandwidth consumption will simultaneously deter customers from using the connection to host a web site or share Napster files, while increasing the revenue collected from those users who persist. Flexible mediation systems can allow service providers to offer free or low-cost usage up to a low threshold, while increasing the price for those users who host home servers. This turns the abusive and costly subscribers into lucrative customers.

D. Excessive Bandwidth use – High Traffic subscribers.

Data collected in production HSD networks has shown that a small percentage (<1%) of subscribers account for over 20% of the total bandwidth consumed. This is typically due to these select few subscribers taking advantage of the flat rate billing to download huge amounts of content from Internet File Servers,

¹⁰ Some subscribers are rebooting the Cable Modem at frequent intervals in order to destroy the MIB Interface counters which are stored in volatile memory.

¹¹ TCP Ports are typically implemented as a 16 bit counter $\sim 2^{16} = 65,536$

or leaving streaming audio (NetRadio) service on for hours, even in their absence. More sophisticated users run 'bots' against newsgroups in order to fill their hard-drive with all of their favorite pictures and software which are automatically downloaded throughout the day and night.

- Who – regular subscribers who fail to realize that NetRadio services consume bandwidth and power users who love newsgroups, etc.
- Impact – Cost, Revenue, QoS
When the HSD system reaches its maximum capacity, subscriber growth is impacted by the limited HFC or peering bandwidth. Dealing with those few (<1%) high traffic subscribers can allow the system to reclaim bandwidth and add as much as 20% more subscribers before upgrading peering links or dedicating more spectrum to HSD services.
- Service Provider Solutions - High traffic subscribers can be addressed in 3 ways

SNMP – collecting MIB information from the CMTS & CM provides the number of packets in / out of a modem since the last modem reboot¹². Collecting this MIB data allows the service provider to identify the top users, but SNMP polls also contribute to the bandwidth congesting the network.

Real-Time Analyzers – Real-time analyzers on the IP infrastructure can produce reports of the top subscribers by total traffic, upstream and downstream traffic. This allows the Service Provider to quickly identify the top users as well as the protocol or services they are using, without adding

¹² The DOCSIS specifications do not require the CMTS to maintain bits in/out counters; these are kept in Volatile memory the CM and are reset to zero following a CM reboot.

bandwidth to the network and regardless of modem reboots.

Usage based billing – Much like how residential long distance and cell phone use is not abused, subscribers are less likely to leave NetRadio services on when they leave the house if they are being charged for the amount of bandwidth or services used.

E. Mis-configured systems – IP theft

Subscribers can mis-configure their client PC either deliberately or erroneously. This usually involves 'IP Theft' where the main PC or a second PC is statically¹³ configured with the wrong IP address. In some systems DHCP servers run in promiscuous¹⁴ mode offering an IP address to every PC the subscriber connects.

- Who – Both Hackers and regular subscribers can decide they want to connect a second or third PC behind the Cable Modem without paying for the additional IP address.
- Impact – Quality of Service
When a subscriber uses the IP address that is assigned to another subscriber, the latter will either experience heavy packet loss or a service outage. Operations staff and CSRs must then spend many hours tracking down the source of the problem if they do not have the proper tools.
- Service Provider Solutions – It is very difficult to identify the subscriber who uses the wrong IP address. In some cases SNMP¹⁵ can be used

¹³ Static configuration is where a subscriber hard-codes the system IP address. This is different from dynamic configuration where the IP address is assigned by a Service Provider maintained DHCP server.

¹⁴ In promiscuous mode a DHCP server does not require a system identifier (such as a host name or MAC address) before it provides an IP address.

¹⁵ This depends on the granularity of the Cable Modem MIB.

to successively check each modem looking to see if its Client PC is configured with the stolen IP address.

F. Attacks

There are a number of IP based attacks for which subscribers can be either the originator or the target. Intrusions (user breaks into an other users system), denial of service^C (TCP synflooding, ICMP attacks, smurfing¹⁶ etc.), virus distribution, portscans¹⁷, forgery campaigns, censorship attempts, etc. are some of the attacks that are seen in IP networks.

- Who – Script Kiddies and Hackers are, by definition, the subscribers who use the Internet to attack other subscribers. These can be random attacks, to learn or test the effects of an exploit, or attacks to gain notoriety in the online community.
- Impact – Quality of Service, Bandwidth, Bad Press
When an HSD network is the recipient, these attacks can consume huge amounts of bandwidth and even generate an outage (such as those which crippled eBay, CNN and other leading web sites in the Spring of 2000).

When an HSD network is the source of an attack, other providers will expect the originating Service Providers operations staff to quickly stop the attack and may not hesitate to shut down peering interfaces in an effort to protect their networks from the attack. When peering links are shut down

¹⁶ Smurfing - Denial-of-service attack consists largely of the use of forged ICMP echo request packets and the direction of packets to IP broadcast addresses rapidly congesting the network and quickly leading to an outage.

¹⁷ In a portscan a user looks at all IP ports to find a vulnerability to exploit perhaps to gain access to, or simply crash the system

quality of service is directly affected as traffic either cannot reach its destination, or must be re-routed over a longer path.

A victim whose computer is attacked may expect the service provider to offer reasonable protection. Even if the service provider doesn't offer a 'managed or firewalled' connection, it may be reasonable to expect that the service provider would identify and block some attack traffic such as preventing the subscriber's entire 512 kbps rate-limited connection from being saturated with attack traffic.

Attack reports are common topics for press articles on 'CyberCrime' or network outages and can damage the brand and the public perception of the HSD service offering.

A Service Provider can discover an attack after being notified by either another service provider, an attacked customer, or multiple subscribers complaining of an outage or poor service or even the press. Once people are complaining about the attack, the service provider is caught in a reactive mode and much of the damage is already done.

○ Service Provider Solutions:

Real-Time Analyzer – A real-time analyzer can monitor IP traffic on the Service Provider's network and alert them to increases in traffic, or certain patterns of traffic such as portscans.

G. Illegal activities

Some behavior is not only against the AUP but is simply illegal. Gambling, child porn, fraud, theft of information, software piracy, copyright infringement etc. are all illegal activities which can be perpetrated with the assistance of a HSD service.

- Who – Fraudulent businesses engage in activities such as gambling, software piracy and fraud. Hackers typically believe software should be free and are known to perform software piracy.
- Impact – While these activities may not affect the HSD Service Provider's network or business model, authorities may require, with proper warrants, the service provider to co-operate in selected investigations by providing data on involved parties who are subscribers. This can consume operational resources as staff parse logs in search of the requested data.
- Service provider Solutions
There is little a service provider can do to prevent or quickly identify this type of illegal activity. The best option is to use a real-time Analyzer system which can track certain activities and store these in a database for later retrieval. However, knowing that service providers did not have the proper tools to track illegal activities the FBI has come up with software called 'Carnivore'¹⁸, which tracks Internet access for law-enforcement.

CONCLUSION

Various forms of abuse in service provider networks can be costly whether in terms of lost revenue, deterioration in quality of service, premature system upgrades, bad press or network downtime. As the subscriber base grows in High-Speed-Data systems, Service Providers are continuously faced with the cost of upgrading capacity. In some cases it may be more economical to reclaim bandwidth lost to abuse than to purchase and upgrade capacity. Since reactive approaches can only attempt to contain damage, reclaiming lost bandwidth is best done with a proactive monitoring approach.

Network equipment such as routers, switches and modems can provide some network information (e.g. SNMP or RMON data) but are fundamentally built to forward packets and provide reliable IP services, and are not designed for the granular real-time data collection required for abuse detection or mediation. Dedicated hardware appliances are the only tools with the flexibility to provide systematic real-time identification of Spammers, home servers, bandwidth 'hogs', IP address theft and some forms of attacks, while operating at wire speeds up to Gigabit Ethernet or OC-12 (both of which are ubiquitous in today's HSD networks).

In addition to identifying various forms of abuse, a scalable software infrastructure (Internet Business Infrastructure) relying on real-time analyzers can help the service provider implement usage based or tiered billing services, which deter from abuse, both across DOCSIS 1.1 services and over the existing legacy subscriber base. Though there are different ways of dealing abusive customers, usage based billing models with granular measurements will both deter abuse and turn those persistent abusers into some of the most lucrative subscribers.

¹⁸<http://www.fbi.gov/programs/carnivore/carnivore.htm>

ABOUT THE AUTHOR

Pat Darisme is a Broadband Strategist and Principal Systems Engineer with NARUS Inc., the leader in Internet Business Infrastructure Solutions. During his time at NARUS he has helped design a decision support and abuse detection system for HSD Service Providers and open access environments. Prior to joining NARUS he was at Excite@Home Network where he worked in Operations and Engineering. He holds a B. Sc. in Electrical Engineering from Universite Laval in Quebec city, Canada. He can be reached at patd@narus.com.

^A Source: Nielsen//NetRatings Global Internet Trends, Q2 2000,

<http://www.eratings.com/news/20000907.htm>

^B Source Forward Concepts, Broadband in the local loop '00, 2000

<http://www.forwardconcepts.com/press26.htm>

^C CERT[®] Advisory CA-1998-01 Smurf IP

Denial-of-Service Attacks

<http://www.cert.org/advisories/CA-1998-01.html>

AN ANALYSIS OF AUDIO FOR DIGITAL CABLE TELEVISION RECOMMENDATIONS FOR THE DIGITAL TRANSITION VIA AUDIO METADATA

Jeffrey C. Riedmiller
Dolby Laboratories Inc.
San Francisco, CA

Abstract

With an increasing number of cable systems deploying large volumes of digital set top boxes and digital services it has become increasingly apparent that large audio level discrepancies can exist between analog and digitally tiered programming. Since the local cable system usually has no control over the digital service audio levels it is up to the programmer(s), in most cases, to rectify any of these discrepancies. Programmers offering digital services are able to take advantage of some very unique tools included with their encoding systems to assist with many of these issues. These tools include the generation of audio metadata that is carried within the coded audio bitstream to the digital subscribers set top decoder. In this paper, the impacts of supplying digital programming with “valid” audio metadata vs. “default” or “invalid” audio metadata from the digital cable subscriber’s perspective will be analyzed. However, the proper use of audio metadata by the programmer also assumes that the millions of deployed set top decoders have been designed with a clear understanding of the entire system as well. Recommendations for digital set top designs are needed for the cable industry and will insure that the effort to generate “valid” metadata by programmers will result in the expected behavior within the digital set tops. To this end, a reference model is also proposed in this paper.

Audio Metadata Origin and Digital Audio Delivery Requirements

The digital audio delivery method currently used by cable networks and cable system operators is the Dolby Digital (AC-3) audio coding system and has been adopted by the Society of Cable Telecommunications Engineers (SCTE) for use in digital cable systems within the United States.

When the Dolby Digital (AC-3) coding system initially developed for cinema applications was adapted into a form suitable for consumer use, the differing and practical needs of consumers were evaluated. As a result, it was decided to create a signal representation syntax that would allow a single encoded bitstream to be decoded into a form useable by nearly every potential listener. Many factors were considered when developing this syntax. The first was how a potential listener may receive the decoded content. In the case of a digital cable set top box this could be via the line level, RF or Digital Interfaces. Second, the number of playback channels in use. Thus allowing multichannel audio programs to satisfy consumers with only 2-channel decoders. Third, the amount of dynamic range desired knowing that consumer product capabilities and individual tastes vary. Fourth, provide a means for level matching between programs and channels. Finally, the allowable dynamic range on a modulated RF interface to control peak program levels similar to present day broadcast practices.

This syntax or “extra” information carried in the encoded bitstream is commonly referred to as metadata.

What is Audio Metadata?

Audio metadata is the data sent in tandem with the coded audio signal to describe this signal to the receiver/decoder. In the case of Dolby Digital (AC-3), audio metadata is used along with subscriber input to control how the audio program is presented in the home, including modification of the program dynamic range or the number of playback channels. For the purposes of this paper, some of the more important metadata parameters and their uses are shown in Table 1 (Table 1 is located on the last page of this paper). These “key” metadata parameters will have the greatest effect on how a listener or subscriber experiences a program. A complete list of transmitted metadata can be found in ATSC Document A/52. It is important to note that these parameters are normally determined during the Dolby Digital encoding process by configuring the encoder. The Informational metadata shown in Table 1 is used by the Dolby Digital decoder to optimize the decoding process by specifying the number of channels used and the program configuration. It determines where to route the decoded audio channels, whether or not the channels should be downmixed (i.e. 5.1 channel to 2 channel), and what type of service is present (i.e. Complete Main, Music & Effects, Dialog, etc).

Level Control metadata sets the overall program playback level for loudness normalization (*dialnorm*). It also determines how various mixdowns from 5.1 channels to one, two, or three channels are achieved (*cmixlev*, *smixlev*).

Dynamic range control metadata has the ability to control the dynamic range of the decoded program at the receiver. The values

are calculated in the encoder based on dynamic range profiles selected by the content creator or the programmer. There are five standard profiles: Film Standard, Film Light, Music Standard, Music Light, and Speech. As noted in Table 1, a *compr* word is generated once per Dolby Digital frame (32 msec) and a *dynrng* word is generated once per block (5.3msec). In addition to generating dynamic range control words for presentation purposes, the encoder also generates control values to prevent clipping in the decoder. This is necessary due to the higher levels produced by combining channels if downmixing is required. The details of Dynamic Range Control will also be discussed later in this paper.

Digital Set Top & Subscriber Requirements

In most cases the digital cable set top box has three primary interfaces for supplying audio to subscriber owned equipment. The two most common interfaces include the ch.3/4 remodulator (via RF output) and the analog baseband L&R line level outputs (via RCA connectors). The third output, if available, gives the subscriber access to the raw Dolby Digital (AC-3) bitstream which can be decoded externally by a home theatre system, if desired. It is also important to note that in two of these cases (i.e. baseband and RF), the decoded digital program must satisfy and provide the digital subscriber a seamless listening experience when switching between digital and analog channels.

At first glance, this level of operability may seem impossible to achieve. However, with a thorough understanding of the Dolby Digital audio delivery system and the proper application of audio metadata these requirements can indeed be met, and therefore increase subscriber satisfaction.

Since Dolby Digital (AC-3) programs need to coexist with analog (NTSC) program sources, the gain structure of the set top box should satisfy the following criteria: 1. All sources whether analog or digital from the set top should provide similar dialog levels at the line level outputs and in any operating mode. 2. In most cases the set top includes a RF modulator and the dialog level from digital sources should match that of typical analog (NTSC) broadcast practices. Note that this can only be accomplished with the Dolby Digital decoder operating in RF Mode. 3. The line level outputs should provide, when connected to the line inputs of a television, similar dialog levels to that of the television's tuner.

The last paragraph referred to the "dialog level" of a particular source, analog or digital. This level is used as a metric to establish a subjective level match for the subscriber from any type of source. This unique approach to level matching may seem quite different at first glance but experience has shown that a very large percentage of television viewers base their listening enjoyment of a particular program on how intelligible the dialog is. Thus, the reproduced dialog level now becomes a reference for the programmer and sound mixer to use to ensure the consistent reproduction of their programming. In the case of motion picture mixing, the dialog loudness is largely adjusted, and thus normalized, by ear in a room similar to a typical theatre with the acoustic "gain" carefully set to a standard value. The motion picture industry has been performing this "dialog normalizing" practice for years to compensate for the inability of theatres to adjust levels on a film-by-film basis.

On the other hand, many television broadcasters still employ the same practices and audio processing techniques that have been around for some time. First and

foremost, the dynamic range of the program is often limited and the level of speech is such that waveform peaks may frequently reach 100% modulation, which, in turn, leaves a very limited amount of headroom for music and sound effects to make a dramatic impact on the listener. However, Hollywood theatrical films offered by premium cable programmers' use the same soundtrack mixed for theatres that have reproduction systems with virtually unrestricted dynamics capabilities. These films often have far greater dynamic range and headroom requirements than typical television programming. Also note that because of this increased headroom, dialog will be placed lower in level within the mix, and will not generally match normal television program dialog levels. This is due to their different peak-to-average ratios. For reference, movies typically have a dialog level of ~ 27dB LAeq below maximum level and typical NTSC broadcasts having a typical dialog level of ~ 17dB LAeq below 100% modulation. (LAeq will covered in the next section)

By now it should be obvious that presenting digital subscribers with these two very different types of programming while simultaneously keeping their equipment from being driven into overload could certainly cause complaints. This could result from the differences in the high peak-to-average ratio (i.e. wide dynamic range) sources such as film and a low peak-to-average ratio (i.e. narrow dynamic range) source such as a news broadcast. Often, a premium cable programmer solves this by processing their audio programming to more closely match current television broadcast practices. By restricting the dynamic range and therefore increasing the average dialog level of the program, sources can be made to match. While this solution is quite acceptable for programming that does not require the dynamic range of typical theatrical releases, it

forces this restriction on all programming and therefore ignores the capabilities of the Dolby Digital system as well as the intentions of the original sound mixer.

With the Dolby Digital system, there are no assumptions about the program levels, so the source can be presented to the encoder with almost any operating practice or philosophy. The system is quite capable of delivering programming with differing average dialog levels and dynamic range without potentially annoying subscribers as they switch through various types of programming that exist on a typical cable system. However, this can only be accomplished when the system is utilized and implemented properly by the programmer, set top manufacturer and in some cases the local cable system. For the programmer itself, it is extremely important to understand the use of Dolby Digital metadata and the parameters that generate and control it in their encoders. One of the most important encoder parameters is the Dialog Normalization value – *dialnorm*. Once this value is entered in the encoder, it is used by the encoder and carried in the bitstream to the decoder as a part of metadata. The *dialnorm* value provides, among other things, the digital subscriber with normalized dialog levels between differing types of digital sources and programs. The next two sections will cover this key metadata parameter in detail.

Level Measurement Practices

Broadcasters usually control their program levels by a volume unit (VU) meter or a peak program meter (PPM) and it is important to note that both are used to read signal voltages, and therefore make no attempt to measure subjective loudness. Thus several different voices, adjusted in level so that they all deflect meters to the same mark, may sound somewhat different in level to the listener. Over the years there have been several attempts to design measurement devices (i.e.

meters) which give the broadcaster results that more closely match subjective loudness. One measure that gives us results that correlate closer with subjective loudness is called the equivalent loudness method (LAeq) which is the long-term average A-weighted level. Leq itself can be defined as the level of a constant sound, which in a given time period has the same energy as a time-varying sound. Matching the levels of different voices so that they give the same LAeq measurement will deliver closer subjective loudness than either PPMs or VU meters. PPM and VU meters are also frequently used to measure and/or align to a pre-determined “house” reference level, and thus only have an arbitrary relationship to the dialog or speech level within a given program. For example, if a VU meter and a PPM meter are calibrated to display a reference tone equally, and speech that averages 0VU is applied to both, the PPM meter will indicate levels considerably above its reference level and possibly above the maximum permitted level. On the other hand, speech that averages at the PPM reference, will most likely indicate many dB below 0VU (our original reference). This confirms the very important idea that the reference level is not the same as the speech or dialog level of a program. Instead, perhaps we should define the relationship between reference levels and the actual dialog (or speech readings) of a given program. With this relationship known, we can then easily describe how to make programming controlled with VUs or PPMs approximately deliver standardized dialog level in Dolby Digital (AC-3).

The standardized dialog level for Dolby Digital (AC-3) is defined as –31dBFS and is quantified using the equivalent loudness method that is A-weighted (LAeq) as stated above. This is the level of spoken dialog present in a single channel, such as the center channel of a 5.1-channel program. For stereo

and downmixed (i.e. 5.1 channel downmixed to 2-channel as in the case of digital cable set top boxes) programs, the level is equivalent to -34dBFS LAeq in both channels simultaneously. This is because dialog will be present in both the left and right channels and will combine acoustically for a ~3dB increase in level.

As stated earlier, dialog levels in conventional analog television audio are roughly normalized, in many cases, by the use of automatic dynamics processing. Recent measurements indicate that they are typically ~ -17dBFS, where 0dBFS would equal the maximum permitted FM carrier deviation (25kHz in the U.S.) More detail on the importance of this will be presented in the following sections.

Dialog Normalization – *dialnorm*

Perhaps the most often misunderstood Dolby Digital (AC-3) metadata parameter is the Dialog Normalization value or *dialnorm* value. The Dialog Normalization parameter describes the long-term average dialog level of the associated program and is specified on an absolute scale that ranges from -1 dBFS (dB Full Scale Digital) to -31 dBFS LAeq.

Encoded Dolby Digital elementary streams carry the audio signals that are fed into the encoder with no changes to the level or the program's dynamics. It is the metadata, and in this particular case the *dialnorm* value, that is used by the decoder to adjust the reproduced level of audio programs which will then reproduce the dialog at a consistent or uniform level. It is also important to note that all decoders are required to make use of this metadata parameter and apply the proper level normalization/attenuation (based on the transmitted *dialnorm* value) to the decoded audio program.

As an example of what the dialog normalization parameter is capable of, imagine switching between a news program and a wide dynamic range movie. In order of magnitude, these items may have average dialog values of about -14 and -28 dBFS LAeq respectively. Thus, if a listener sets the playback level (using a volume control) to comfortably reproduce the news program and then switches to the movie, its dialog will be about 14 dB (28-14) quieter than the voice of the newsreader, and probably unintelligible. This will likely force the listener to adjust the volume control every time they switch between the two. In this example, note that the quieter source is the movie dialog. The reason for this as stated earlier is movie mixers use a standardized acoustic level for dialog while mixing. With digital formats this is equivalent to between -25 and -31 dBFS LAeq. Since movies constitute a large portion of the material to be conveyed by Dolby Digital, this standardization is retained. This means that *all* dialog (movies, sports, news, etc) should emerge from a Dolby Digital decoder at about -31 dBFS LAeq. However, for this to be the case the programmer **must** properly set the *dialnorm* value in the Dolby Digital encoder. In the example above, *dialnorm* needs to command 17 dB of attenuation on the news program, and only 3 dB on the movie to create an acceptable level match between the two upon decoding. In other words, *dialnorm* in a sense acts as an automatic volume control. In the case of programming that does not contain dialog, such as music, *dialnorm* should be adjusted appropriately to match programming that includes spoken dialog.

Remember that the *dialnorm* value itself is **not** the standard operating level of a facility (i.e. -20dBFS = + 4dBu or 0VU) but the equivalent of the spoken (not shouting or whispering) dialog level with respect to digital full-scale. The attenuation introduced

in the decoder is $(31 + \text{dialnorm value})$ with the *dialnorm* value being negative. Speech in a single channel (i.e. center channel of a 5.1 channel program), with an LAeq value of -31 dBFS should have -31 entered in the encoder, which commands 0 dB of attenuation in the decoder. If speech is present in two channels simultaneously (a stereo program) with an LAeq value of -34dBFS in each channel, -31 should be entered in the encoder. Similarly, the news bulletin (mono) with speech at -14 dBFS LAeq requires a *dialnorm* setting of -14, giving $31 + (-14) = 17$ dB of attenuation so that speech from the newsreader comes out of the decoder at -31 dBFS. If the source material is recorded at a lower level resulting in peaks that do not approach digital full-scale, less attenuation is needed, and the *dialnorm* value moves closer to its minimum value of -31. Thus, if the news bulletin used in the example above had a dialog level of -20 dBFS LAeq rather than -14 at its source, a *dialnorm* setting of -20 would yield the standard level (-31 dBFS) at the decoder outputs while operating in Line Mode.

This unique feature alone has the potential to minimize the average channel to channel dialog level problems that have plagued cable television for years and subsequently has forced many premium service programmers to process their audio to match the typical dialog levels and dynamic range of other non-premium programming. The Dialog Normalization parameter enables the premium service programmer to achieve an acceptable program to program level match while minimizing the need for the traditional and irreversible heavy processing that is commonly applied to most non-premium type programming. This leaves ample headroom for feature films regularly offered by premium network providers while simultaneously maintaining an acceptable level match with typical programming that has limited dynamic range. If the entire

delivery system is to function as intended, an understanding of set top decoder operating modes is required. It is also critical to ensure that the decoder is operating in the correct mode, which can differ in certain applications. Unbeknownst to many cable programmers and local cable system technical personnel, the digital set top that their particular system may deploy can potentially offer and/or default to different operating modes which impact the subscriber in different ways. These modes may disable the ability to offer the subscriber an acceptable level match between analog and digital programming.

Digital Set Top Decoder Operation

This section defines the importance of digital set top decoder operating modes and how each of them are applied to allow us to satisfy our three original requirements which are restated here:

Digital Set Top Output Level Requirements:

1. All sources whether analog or digital from the set top should provide similar dialog levels at the line level outputs and in any operating mode.
2. In most cases the set top includes an RF modulator and the dialog level from digital sources should match that of typical analog (NTSC) broadcast practices. Note this can only be accomplished with the Dolby Digital decoder operating in RF Mode.
3. The line level outputs should provide, when connected to the line inputs of a television, similar, dialog levels to that of the television's tuner.

Dolby Digital decoders found in consumer products, in general, can operate in two modes. Each of these modes has a specific application, and care **must** be taken when the set top is designed and deployed to insure that the intended mode is used by default.

Decoders can be found in many places including Digital Cable set top terminals, Home Theatre systems, consumer satellite Integrated Receiver Decoders (IRD), DBS receivers and Commercial IRDs that are used in cable headends, and Cable turnaround uplink facilities. It is important to note that the default operating modes in each of these cases may vary and are based on that particular device's function within a given system.

Line Mode Operation:

Line Mode operation generally applies to the baseband line level outputs from two-channel set-top decoders, two-channel digital televisions and multichannel Home Theatre decoders. It is important to note that Line Mode operation is a requirement for all digital cable set top boxes that have analog baseband (line level) outputs. With respect to consumer type applications, a decoder's outputs operating in this mode will typically be connected to a much higher quality sound reproduction system than that found in a typical television set. In this mode, dialog normalization is enabled and applied in the decoder at all times. Further, in this mode the normalized level of dialog is reproduced at a level of -31 dBFS LAeq, but **ONLY** when the transmitted *dialnorm* value has been correctly adjusted for a particular program. In general, with the reproduced dialog level at -31dBFS, this mode allows wide dynamic range programming to be reproduced without any peak limiting and/or compression applied as may be intended by the original program producers. Further, since the Dolby Digital (AC-3) digital audio coding system can provide more than 100dB of dynamic range, there is no technical reason to encode the dialog at or near 100% as is commonly practiced in analog television systems. This allows the delivery system to meet one of its goals of being able to deliver high impact

cinema type sound to the digital subscriber's living room.

RF Mode Operation:

RF Mode is intended for products such as cable and satellite set-top terminals that generate a monophonic and/or downmixed signal for transmission via the channel 3/4 remodulator that feeds the RF (antenna) input of a television set. This mode was specifically designed to match the average reproduced dialog level and dynamic range of digital sources to those of existing analog sources such as NTSC and analog cable TV broadcasts. In this operating mode dialog normalization is enabled and applied in the decoder at all times. However, the dialog level in this mode is reproduced at a level of -20 dBFS LAeq **ONLY** when the transmitted *dialnorm* value is valid for a particular program. The Dolby Digital decoder introduces an +11 dB gain shift and thus the maximum possible peak to dialog level ratio is reduced by 11dB. This is achieved by compression and limiting internal to the decoder. It is important to note that digital set top boxes which include an RF modulator are required to provide RF Mode in addition to Line Mode operation and also **must** include some way of changing operating modes.

Figure 1 compares the signal relationships in the decoder for both Line and RF operating modes. Notice the reproduced dialog level and dynamic range of each mode.

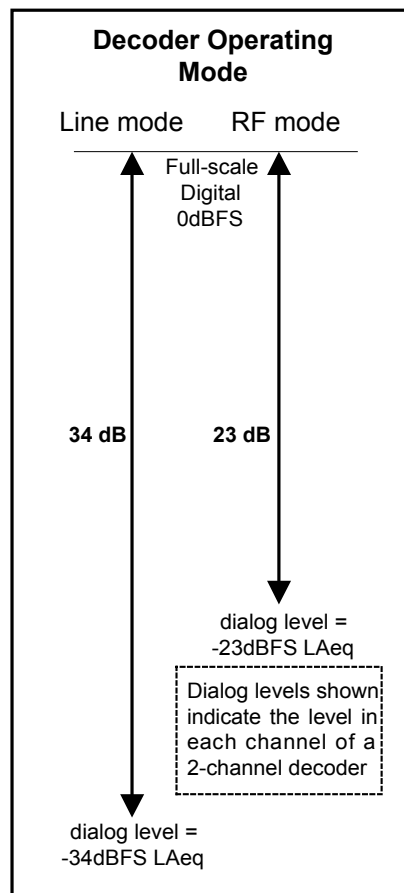


Figure 1

Decoder Operating Mode Levels - Line vs. RF

Digital Set Top Gain Structure

The digital cable set top is required to provide the subscriber with both analog and digital programming. It is imperative that digital set top manufacturers fully understand both decoder operating modes as well their applications within a given set top design.

The audio subsystems within the set top include:

1. NTSC Tuner

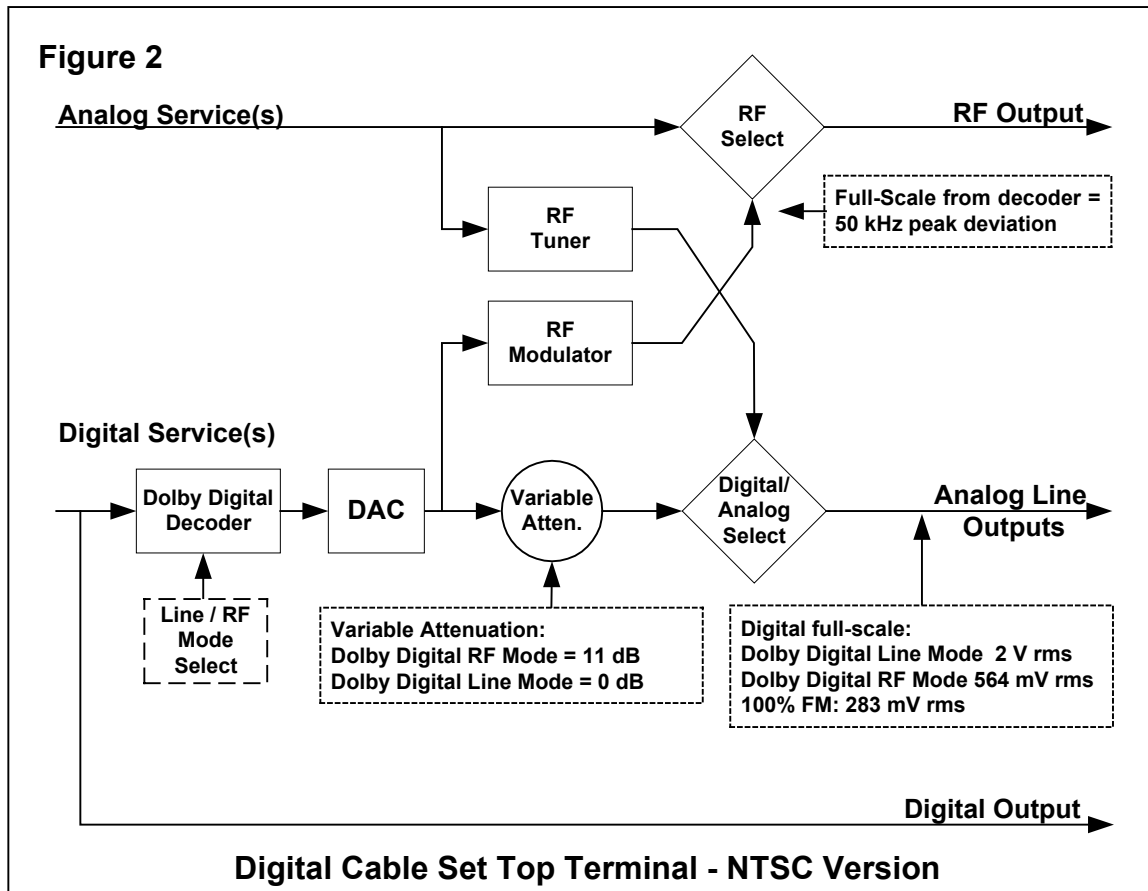
2. 2-channel Dolby Digital decoder IC and associated digital to analog converter
3. Channel 3/4 FM modulator

Beginning with the FM modulator, we shall define 0dBFM as the maximum FM deviation (25kHz) for NTSC broadcasts. Both theory and recent measurements indicate that NTSC tuners have ~ 6dB of headroom above 25kHz deviation while receiving a mono signal. As will be seen, this works to our advantage.

With a Dolby Digital decoder operating in RF Mode, the normalized dialog level is a -20dBFS LAeq (-23dBFS LAeq in each channel), since most digital cable set tops utilize a 2-channel decoder. In some set top designs the output of the decoder IC may feed digital multipliers, used for digital control of the volume before being presented to the digital-to-analog converters (DACs). With this in mind, the gain structure should be provisioned initially with volume controls set to unity gain.

To set the proper amount of gain into the modulator for digital sources, the two analog signals from the DACs should be combined and fed to the mono RF modulator with gain such that a correlated tone (i.e. 400Hz) in both channels (L&R) at 0dBFS yields +6dBFM (200% modulation). When the set top is tuned to a digital source AND this particular source has the correct dialog normalization value associated with it, the dialog should emerge from the decoder operating in RF Mode at -23dBFS LAeq in each channel. Referring to Figure 3, it can be seen that with the RF modulator gain structure set as indicated above, the -23dBFS dialog becomes -17dBFM and matches the typical dialog level of NTSC broadcasts. On the contrary if the Dolby Digital decoder is operating in Line Mode the decoded dialog level for the same program must be 11dB lower.

Figure 2



Recall that in Line Mode peaks can be 11dB higher with respect to the dialog level and in order to avoid overload the average level must then be 11dB lower

Therefore, the normalized dialog level in Line Mode being -34dBFS LAeq in each channel, then becomes -28dBFM . Obviously, under these conditions the digital subscriber will experience a severe level mismatch between analog and digital sources. Hence the requirement, as stated earlier: digital set tops which include an RF output **must** also include the capability of placing the internal Dolby Digital decoder into RF Mode. Note the dialog levels for both operating modes in Figure 3.

The line level outputs on digital cable set top boxes will, in most cases, feed the line inputs of a stereo television, VCR or a home theatre system. And since it may be desirable for set top output levels to match the standard output levels of other consumer audio equipment such as CD and DVD players we must consider the following. Many consumer audio products (with line level inputs) are usually designed to accept a level of 2Vrms maximum and not much more. Therefore the digital set top box must be designed such that digital full-scale signals are able to deliver 2Vrms at the line level outputs.

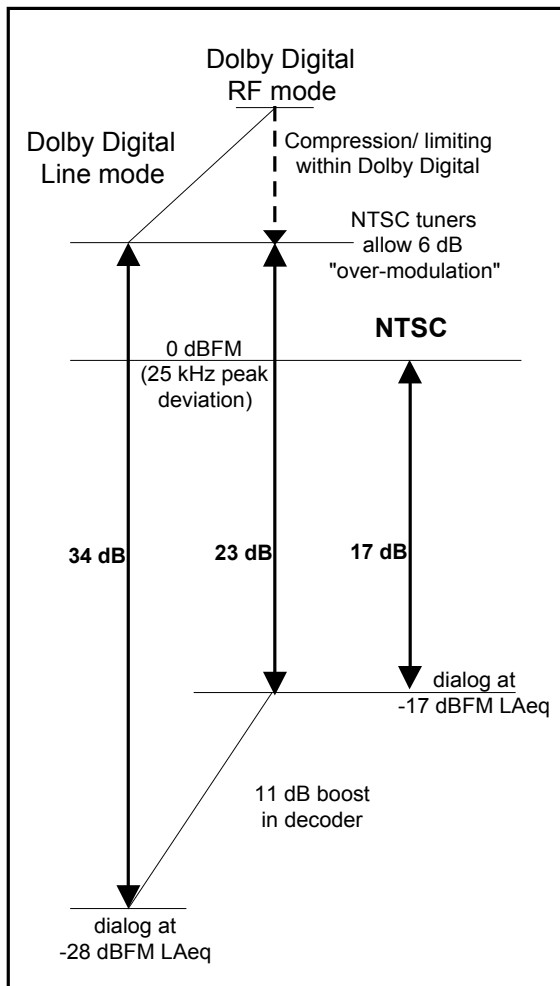


Figure 3

Set Top Gain Structure

To satisfy one of our earlier requirements of having matched dialog levels in both decoder operating modes at the line level outputs, we shall re-state that full-scale (0dBFS) signals from a Dolby Digital decoder operating in Line Mode should give 2Vrms at the line output connectors. Hence, the standard normalized dialog level (quantified using the LAeq method) in Line Mode will be 34dB below 2Vrms and equates to 40mV rms in each channel. If the decoder is deliberately set or defaults to operating in RF Mode, it becomes necessary to remove the 11dB of boost applied in RF Mode before applying the

signal to the line level outputs **ONLY**. See figure 2 & 4. This will then produce the identical LAeq dialog level at the line level outputs of the digital set top box as when the decoder is in Line Mode (40mV rms).

For analog services (NTSC) delivered to the digital set top, our gain structure assumes that the normal dialog level is about -17dBFS. If the NTSC dialog level is to match digital sources, the maximum deviation of 0dBFS should be 17dB above the normalized dialog level of the Dolby Digital decoder (40mVrms). Therefore, 17dB above 40mVrms results in 0dBFS being equal to 283mVrms. See Figures 2 and 4 for more information.

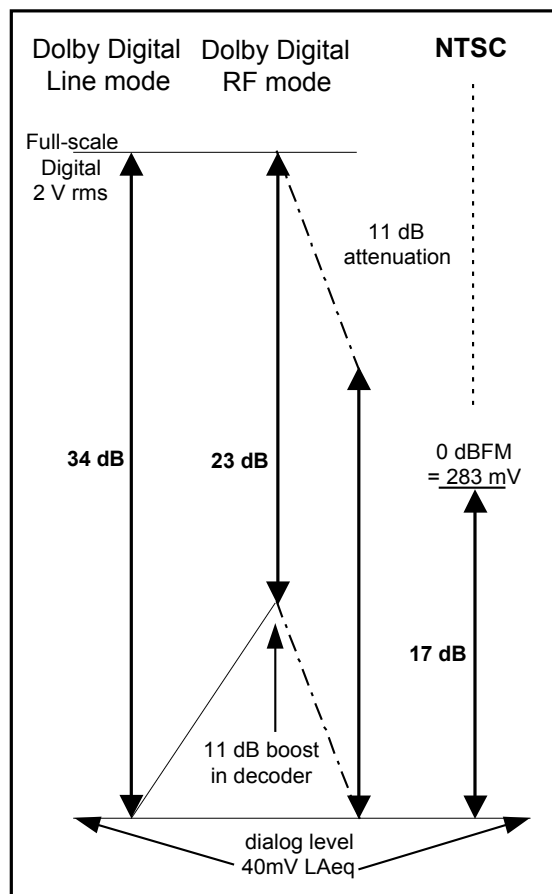


Figure 4

Set Top Box Line Output Levels

Dynamic Range Control

The Dolby Digital system conveys audio without altering its dynamics. Unlike any previous broadcast delivery mechanism, it gives the listener the option to hear the program as the original program producer intended, even if that means that it ranges from barely audible to extremely loud (i.e., wide dynamic range).

Typical analog broadcast processing usually forces the audio program level towards full modulation for a substantial portion of the time, lowering the peak-to-average ratio thereby eliminating most of the dynamic range. A benefit that comes with this is an approximate normalization of the average listening level. In other words, the same device that reduces the dynamic range determines the average volume. With Dolby Digital there are no technical pressures to reduce dynamic range, and the average volume is addressed by the *dialnorm* parameter discussed earlier. Further, the need for dynamic range compression can be considered to be independent of average listening levels.

The Dolby Digital system utilizes a unique approach to applying dynamic range compression to audio program material. The purpose of the algorithm itself is to make the decoded audio levels closer to the dialog normalization level, amplifying material that is lower in level and attenuating material that is above the dialog level. Rather than reducing the dynamic range in a non-reversible way, the Dolby Digital encoder generates compression gain words that are carried along in the Dolby Digital bitstream as part of metadata. These gain words are calculated based on a number of separate input parameters, including the selected dynamic range compression preset (see the compression profile section), the level of the program material itself, and the program

dialog normalization value. Hence, setting the *dialnorm* value properly is a critical first step in calibrating the dynamic range compression system.

The DRC (Dynamic Range Control) algorithm supplies two types of gain words to the decoder, *dynrng* and *compr*. The *dynrng* variable is applied to the decoded audio when the decoder is operating in Line Mode. This control signal is generated based on the original program producer's artistic choice of a compression profile preset, which is included in the Dolby Digital encoder. Since many of consumer decoders default to Line Mode, the programmer may choose, via the compression profile preset, the one that suits the needs of most of their audience. Individual listeners may have the option to choose to decode the program with all of its original dynamic range depending on the type of decoder they have. However, studies and experience have shown that in most cases the majority of television viewers would never want or need the full dynamic range of the audio that this system is capable of delivering.

The *compr* variable is applied to the decoded audio when the decoder is operating in RF Mode. This control signal is generated to insure that the peak modulation of the RF remodulator is controlled to an acceptable value when the dialog is modulated to a value that is similar to current analog NTSC practices. This provides the listener who chooses to use the RF output of a set top terminal with acceptable results without affecting listeners that may be using the line level outputs and/or an external Dolby Digital decoder in a home theatre system.

Downmixing, Overload Protection and Dynamic Range Control

When a multichannel program is downmixed within a Dolby Digital decoder, the downmix

coefficients are generally fixed values. Downmixing is performed in the digital domain (except in the case of a 2-channel to mono downmix), and obviously there is the possibility that the downmix will overload/clip the digital-to-analog converters (DACs) in the decoder. If the fixed downmix coefficients in the decoder were chosen so that downmixing a multichannel program would never overload the DACs, many of these downmixed programs would sound quieter than the same program reproduced in a multichannel mode or a mono or stereo program that did not require downmixing. With this potential problem in mind, the actual downmix coefficients were chosen to give a more satisfactory match in output level between downmixed and non-downmixed programs. However, the caveat with this choice is that it might result in output overload while downmixing on the rare occasions that a multichannel program approached digital full-scale on all channels simultaneously. Remember that most programming typically demands level attenuation within the decoder via the indicated *dialnorm* value. This attenuation is applied in the digital domain prior to downmixing, and in turn reduces the probability of overload the DACs. However, the 11 dB of boost provided by the decoder in RF Mode does once again increase the probability of overloading/clipping the DACs while downmixing. These different situations were taken into account and are addressed within the dynamic range compression algorithm.

To assist the system in predicting a possible overload condition the encoder generates several possible downmixes in parallel with the computations taking place that generate *dynrng* and *compr* gain words. These downmixes are used to estimate the worst-case peak level at a decoders output while taking into account the value of *dialnorm* that

is currently being used. This separate process then calculates the gain reduction needed to prevent an overload condition for decoders operating in both Line & RF Modes. The two values of gain reduction derived from this process, one for Line Mode and one for RF Mode, are compared to the gain reduction values demanded by the selected dynamic range compression profile. The larger value is then substituted for the compression figure in the *dynrng* and *compr* control word. With typical multichannel programming, protection limiting is rarely needed or necessary when the accompanying *dialnorm* value is correct. As an example, imagine the highly unlikely case where a five-channel program reached full-scale simultaneously in all channels. This represents a sound roughly 30 dB higher than standard dialog level. In this case, downmixing with the decoder operating in Line Mode and preventing an overload would require that a total of 11 dB of gain reduction be applied from *dialnorm* and *dynrng*. In other words, the attenuation requested by both the *dialnorm* value and the calculated *dynrng* gain word (actual value based on compression profile selected in the encoder) must be 11 dB or more to keep heavy protection limiting from taking place. In RF Mode, this condition could require 22 dB of gain reduction (*dialnorm* + *compr*). This may be a good reason to always use an artistic compression profile in the encoder when given the choice.

When the bitstream is decoded, the decoder itself is responsible for applying these gain words to the reproduced audio. The particular gain words used are determined by the operating mode that the decoder is in. The decoder can be instructed to provide the full amount of compression indicated by the gain words, reduce or scale the amount of compression applied, or even apply no compression at all (as long as the decoder is operating in Line Mode and not performing a

downmix). This allows the end user to adjust the amount of dynamic range compression applied based on their individual taste.

Compression Profiles & Presets

As stated earlier, in many applications Dolby Digital (AC-3) will be compared and heard along side analog terrestrial and cable programming sources. There are several Dynamic Range Compression profile presets included in the current Dolby Digital encoder that permit degrees of processing that are similar to present day broadcast practices. They all include a null band, which is a region in the middle of the dynamic range where the gain is fixed at unity. (i.e. no boost or cut) The ends of this null band are referenced to the value of *dialnorm* so that dialog or average program lies within the null band and is not subject to gain variation. Applying dynamic range compression in this way provides two important advantages. First, program material with an already restricted dynamic range will generally remain within the null band and in turn, is not subjected to any further dynamic range compression, and second, dialog does not modulate background noises. The result is that this type of processing reduces the dynamic range without producing audible side effects such as transient distortion or gain pumping often associated with broadcast processors. Figure 5 is a graphical representation of unity gain within the compression profile null band.

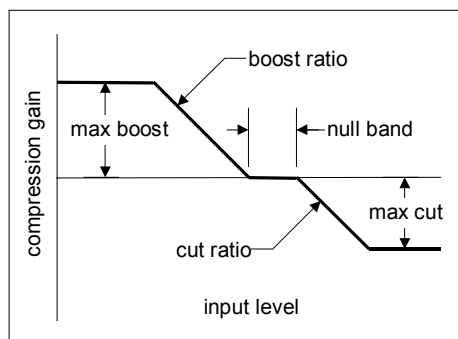


Figure 5 – Example Dynamic Range Profile

It is assumed that a typical wide dynamic range program, such as a film, has peak levels at least 10 to 15dB above the dialog peak levels and will probably be unacceptably loud for many subscribers. Thus, if the dialog level lies within the null zone, typically 10dB wide, the only way to restrict these changes in level is to use a high compression ratio above the top end of the null zone. Below the null zone, the maximum amount of gain to be applied is cause for concern. In film, raising low-level sounds excessively would sometimes reveal unwanted background disturbances, such as camera and traffic noise, which the original program producer did not intend to be audible. Since the sounds intended to be heard are rarely much quieter than dialog, film soundtracks would only require small degrees of low-level boost. Thus the low level compression ratio is not critical and limiting the maximum gain for film soundtracks to only about 6dB will provide acceptable results.

Currently, there are five Dynamic Range Compression profiles in addition to None included in the Dolby Digital Encoder. These profiles are: Film Standard, Film Light, Music Standard, Music Light and Speech. For more information on the specifics of each compression profile please consult the Dolby Digital Professional Encoding Guidelines Manual.

Also note, if the “None” compression profile is selected, the dynamic range compression algorithm generates the desired compression gain words set to 0 dB gain (i.e., no cut or boost). However, the “None” profile does not disable overload protection and it is possible that the actual compression gain words in the Dolby Digital bitstream will be less than 0 dB.

Current Observations and Practices

A recent analysis has indicated that much of the digitally distributed cable programming is encoded with incorrect *dialnorm* values and in many cases is utilizing the preset default values that many Dolby and Dolby licensed encoder manufacturers have implemented. This most likely stems from confusion over the purpose of *dialnorm* and how to properly set it for a given program. Note that default values may vary from manufacturer to manufacturer, further aggravating the problem. The documentation that usually accompanies these very complex encoding systems typically does not include clear explanations on how *dialnorm* and other metadata parameters inter-relate, how to set them, and in the end, how they affect the listeners and subscribers. At least one manual simply states that setting the *dialnorm* value to -31 will cause no attenuation to take place in the decoder. While this statement is true, it should be realized by this point in the paper some of the serious drawbacks of doing so.

Further complicating things, some of these potential drawbacks become increasingly worse if the correct *dialnorm* value for a given program is a greater distance (with respect to dialog level) from a typical default *dialnorm* value of -31. Take the case of a program that regularly reaches 0 dBFS and has a very limited dynamic range such as rock music. Encoding this type of program with a *dialnorm* value of -31 appears logical, due to the fact that no deliberate attenuation will take place in a decoder operating in Line Mode. Unfortunately it will, more than likely, cause audible side effects with decoders operating in RF Mode due to protection limiting, and may lead to subscriber complaints. Since a large percentage of digital cable subscribers are utilizing the channel 3/4 RF output of their digital set top terminal to feed their televisions, properly setting the Dolby Digital metadata is extremely important.

Recall that a Dolby Digital decoder operating in RF Mode reproduces the dialog level (based on a valid transmitted *dialnorm* value) at -20 dBFS LAeq which is an +11dB increase in dialog level when compared to a decoder operating in Line Mode of -31dBFS LAeq. For this example, when the encoder calculates and creates the *compr* gain words (see the earlier section on Overload Protection) used by a decoder operating in RF Mode, they are given the responsibility for at least 11dB of overload protection limiting to keep the set top decoder DACs from clipping. This type of protection limiting can be aggressive at times, especially in the case where an artistic dynamic range compression profile was not chosen in the encoder. Had the *dialnorm* value been set correctly in this case, there would be minimal or no overload/heavy protection limiting applied in a decoder operating in RF Mode.

While the correct setting of *dialnorm* for this program will minimize any side effects from heavy protection limiting with decoders operating in RF Mode, it will also lower the reproduced level for decoders operating in Line Mode. This behavior is expected. In the case of popular music productions, with recorded levels frequently approaching 0 dBFS and average levels not too far below that, some attenuation in the decoder is required to better match the average levels found in typical analog and digital broadcast programming.

Recently, many premium service cable networks have installed systems that enable them to offer multichannel audio to their subscribers. With some of these systems it is now possible to easily carry all of the proper audio metadata generated in production, through the entire distribution system, and eventually to the subscriber's digital set top terminal and/or decoders placed in a cable

headend or turnaround uplink. When a service such as this is directly compared to another programmer's service, which is using default and most likely an incorrect *dialnorm* value, it can potentially further aggravate the level discrepancies. In the worst case, this scenario applies to cable systems that are supplying digital subscribers with programming with valid metadata (in this case a correct *dialnorm* value) and simultaneously offering other programming without valid metadata, all in the same channel line up delivered to the subscriber's set top terminal.

For example, a movie typically has a *dialnorm* value of -27 dBFS LAeq. Assuming that the Dolby Digital decoder is operating in Line Mode (dialog reproduced at -31 dBFS LAeq) a program having a *dialnorm* value of -27 (movie) would require 4 dB of attenuation in the decoder to "normalize" the dialog level. In contrast, programming with an actual *dialog* level of -21 and a transmitted *dialnorm* value incorrectly set at -31 would request no attenuation or normalization take place in the decoder. The net result of this is that the two programs are now 10 dB apart upon decoding. This is a good example of how incorrectly setting *dialnorm* on only a few services can affect the perception of those services that have it set correctly. Since some programmers, especially those dealing with legacy content, may never have the need to offer their subscribers multi-channel audio and/or have a system in place to correctly set the audio metadata (in this case *dialnorm*) on a program by program basis configuring *dialnorm* appropriately becomes a challenge.

For this particular broadcaster, if most of their source material is already uniform in dialog level, possibly through the use of processing, but not at the normalized level of Dolby Digital then a "static" or fixed *dialnorm* value

may be appropriate. In other words, utilizing *dialnorm* to introduce a fixed offset in the decoder to make up for the offset between the plants operating level and the normalized Dolby Digital level. Implementing a fixed *dialnorm* value, in this case, still allows the programmer the ability to easily make minor adjustments if needed. This practice of setting *dialnorm* to a valid "static" value has been implemented on several services and has been quite successful in providing much closer channel-to-channel level matches while avoiding overload/protection limiting distortions. A critical benefit is that all of this could take place without the need for the programmer to significantly modify their internal operating procedures and practices.

Analog Tiered Subscribers

As stated earlier, many cable programmers are beginning to use digital compression methods to transport their signals to both analog and digital cable systems. The main difference is where the digital content is decoded. For an analog system or tier, the service is decoded to baseband analog in the cable system's headend facility using a commercial Integrated Receiver Decoder – IRD. The decoded signal in analog form is then modulated and combined with other services/channels onto the cable plant for delivery to the subscriber. Since all of the subscribers on this type of cable system (analog) will receive the signals via the RF Tuner input on their television or set top terminal, typical wide dynamic range movies from a premium service with the headend IRD operating in Line Mode may cause subscriber annoyance when comparing this level to non-premium programming. In this application, Line Mode may offer greater dynamic range than what the analog subscriber would need or want. This is due to the differences in operating practices with various programmers with respect to the average dialog levels below full-scale digital.

A more appropriate IRD operating mode for this application (if available) may be RF Mode compression but without the 11dB boost in overall level reflected at the commercial IRDs outputs. This can be accomplished with -11dB of attenuation placed in series with the line level outputs when the IRD is operating in this condition. This mode limits the dynamic range to that which is typically found in traditional analog NTSC programming. This will give the headend technician a better chance at fine-tuning the audio modulator for this service to match that of other processed or restricted dynamic range programming without gross over-modulation. If RF Mode is not available in a particular IRD, the use of Line Mode and applying *dynrng* scale factors of 1.0 for both cut and boost may also provide satisfactory results in many cases.

Set up and calibration of the downstream analog modulator in the cable headend is similar to the set top box operating in RF Mode example covered earlier. For a monophonic analog modulator, two analog signals from the IRD DACs should be combined and fed to the modulator with gain such that identical (i.e. correlated) tones (400Hz) present in both channels of the decoder at 0dBFS yields +6dBm (200% modulation) as in our set top example used earlier. With the gain set in this manner, programming from the IRD will more closely match other analog services within the cable plant in both dialog level below maximum and program dynamics.

As a side note, many programmers who have begun to include “valid” metadata have witnessed, depending on the material, an attenuation in the decoded level when compared to the level that was presented to the encoder(s). This behavior is normal and is due to the validity of the *dialnorm* value carried in the audio program and Dialog

Normalization always being enabled in the decoder. To explain this, keep in mind that previously many programmers may have been encoding all of their programs with a “default” and perhaps an “invalid” *dialnorm* value of -31dBFS LAeq, which required no attenuation take place upon decoding.

Turnaround Uplink Facilities

Another specific decoder operating mode sometimes implemented in professional and commercial IRDs that is useful in some cable headends or a turnaround uplink is Custom Mode. This mode, among other things, allows the professional user to disable dialog normalization in the decoder. Disabling dialog normalization in a turnaround uplink may not offer any advantages since this type of facility will most likely re-encode the audio programming with locally generated Dolby Digital metadata that may or may not be correct. Consider a scenario in which a programmer is correctly setting all of the Dolby Digital encoding parameters, including *dialnorm* at their uplink. This programming is then distributed to both cable headends and turnaround uplinks. Enabling dialog normalization in the IRD at the turnaround uplink will normalize the incoming average dialog levels to -31dB LAeq. Since the incoming programming is now normalized with respect to dialog level, the downstream uplink encoder can employ a *dialnorm* value of -31 (which in many cases is default) and the control over decoded audio levels is determined at the original programmers uplink encoder. All of this assumes that the turnaround uplink facility does not process the incoming audio before re-encoding and that the programmer has set *dialnorm* properly. Consider the same situation described above but with dialog normalization disabled at the turnaround IRD. In this mode (Custom Mode), the expected dialog level normalization will not take place in the IRD because the *dialnorm* value carried in the

bitstream is ignored. In turn, any differences in dialog level will be distributed to digital subscribers that exist on the downstream side of the turnaround facility, potentially annoying them. Another potential issue arises when the *dialnorm* value used in the turnaround uplink encoder is set to -31: in some cases presenting un-corrected audio levels (i.e. not at -31dB LAeq) may result in unnecessary limiting when the digital set top is operating in RF Mode.

Another very important point with regard to Custom Mode is that when a decoder is operating in this mode and is performing a downmix (i.e. 5.1 ch. > 2 ch.), 11dB of attenuation is automatically applied. This is due to the possibility of disabling *dynrng* compression and dialog normalization, which would increase the potential for overload during downmix condition. In most cases, the attenuation is compensated for externally in the analog domain, so this behavior is transparent to the user. However, it is wise to verify that this is in fact the case.

As a final note remember that most, if not all broadcasters airing a multichannel audio program will most likely have to replace their existing 2-channel or mono audio program rather than sending both simultaneously. This single multichannel audio program then has to satisfy everyone's listening needs. In most cases, a multichannel audio program will need to be downmixed to stereo or mono not only at the subscriber's digital set top but also in a headend or turnaround facility within the commercial IRD as well. Due to the vast deployment of commercial IRDs that cable networks may have, it becomes imperative that they are all operating in the intended modes as defined by the programmer. Many cable networks may not realize that their decoders are operating in inappropriate modes until the multichannel programming actually starts. It is important to note that this change

in audio can affect affiliate cable systems and subscribers differently in each of these cases.

Recommended practices and procedures for commercial decoders are needed throughout this industry and will give programmers and subscribers predictable and consistent results for every type of application.

Conclusion

This paper has begun to outline the issues and opportunities surrounding the deployment of digital programming with an emphasis on the uses of Dolby Digital audio metadata to assist the cable industry in maintaining consistent levels. We have specifically focused on the Dialog Normalization parameter since it has the largest impact on the subscriber. Also provided was an analysis of a reference digital set top box design and how a proper implementation is extremely important. With a better understanding of the overall system, programmers can offer better service to subscribers with the simultaneous benefits of fewer complaints, increased satisfaction, increased customer retention, and best of all, an expanding customer base due to the increased quality of the delivered products.

References

1. Advanced Television Standards Committee A/52, *Digital Audio Compression Standard (AC-3)*, www.atsc.org
2. Society of Cable Television Engineers DVS018, *ATSC Digital Television Standard, Annex B*.
3. Cossette, Stan G & Guillen, Newton G – *New Techniques for Audio Metadata Use and Distribution*, AES Preprint 5044
4. Dolby Digital Professional Encoding Guidelines, Dolby Laboratories Inc.
5. Robinson, Charles Q & Gundry, Kenneth – *Dynamic Range Control via Metadata*, AES Preprint 5028
6. Riedmiller, Jeffrey & du Breuil, Tom

- Audio in Cable, Broadcast, and Satellite
Distribution. Issues & Solutions for the
Digital Transition via Audio Metadata

Acknowledgements

The author would like to acknowledge and thank Ken Gundry, Tim Carroll, Charles Robinson, Stan Cossette, Oren Williams, Craig Todd, Steve Lyman, Tom du Breuil, Dan Groustra, Ned Mountain, Mike Aloisi, and David Sedacca for their support.

Author contact information

Jeffrey C. Riedmiller
Dolby Laboratories Inc.
100 Potrero Ave.
San Francisco, CA. 94103
415.645.5000

Parameter Name	Description	Change Rate	Purpose
acmod	audio coding mode	Program	Informational
bsmod	bit stream mode	Program	Informational
dsurmod	Dolby surround mode	Program	Informational
lfeon	low-frequency effects channel indication	Program	Informational
<i>dialnorm</i>	dialog normalization	Program	Level Control
compr ("RF" profile)	compression gain word	Frame	Dynamic Range Control
dynrng ("Line" profile)	high rate compression gain word	Block	Dynamic Range Control
cmixlev	center mix level	Program	Level Control
surmixlev	surround mix level	Program	Level Control

Table 1 - Dolby Digital Metadata Parameters

AN ENHANCED BROADCAST CHAIN INTEGRATION PROTOCOL

Jay Weber, Ceri Jerome, Prasad Panchangam, Daniel Salo
RespondTV, Inc.

Abstract

Operators and Programmers each have important roles in the delivery of enhanced content, and these roles need technical integration to make ETV work. In particular, programmers are in a position to determine the scheduling of enhanced content and other interactions with programming, and operators are in a position to determine bandwidth provisioning and platform-specific encoding. This paper describes an architecture and a protocol that supports integration of these control points along the broadcast chain.

INTRODUCTION

In this paper, *Enhanced Television (ETV)* is the combination of broadcast programming and synchronized interactive graphics, text, and forms. This category covers such compelling applications as direct response from advertisements, engaging product placements, audience participation, and rich T-commerce. Thus as ETV is a synchronization of programming and interactivity, it requires technical integration of the broadcast signal and the interactive capabilities of distribution and in-home equipment, in particular between programmers and operators.

Many pilots, including some of ours, have assumed that a programmer does all preparation of the signal and the Operator serves as a pass-through, perhaps with signal format conversion. We have found that this approach does not scale for the following two reasons.

First, Operators are deploying a variety of set-tops that have a diverse set of

capabilities and a combination of standardized and proprietary interfaces. Programmers, who generally strive for maximum distribution, cannot prepare the interactive signal components in one way that fits all.

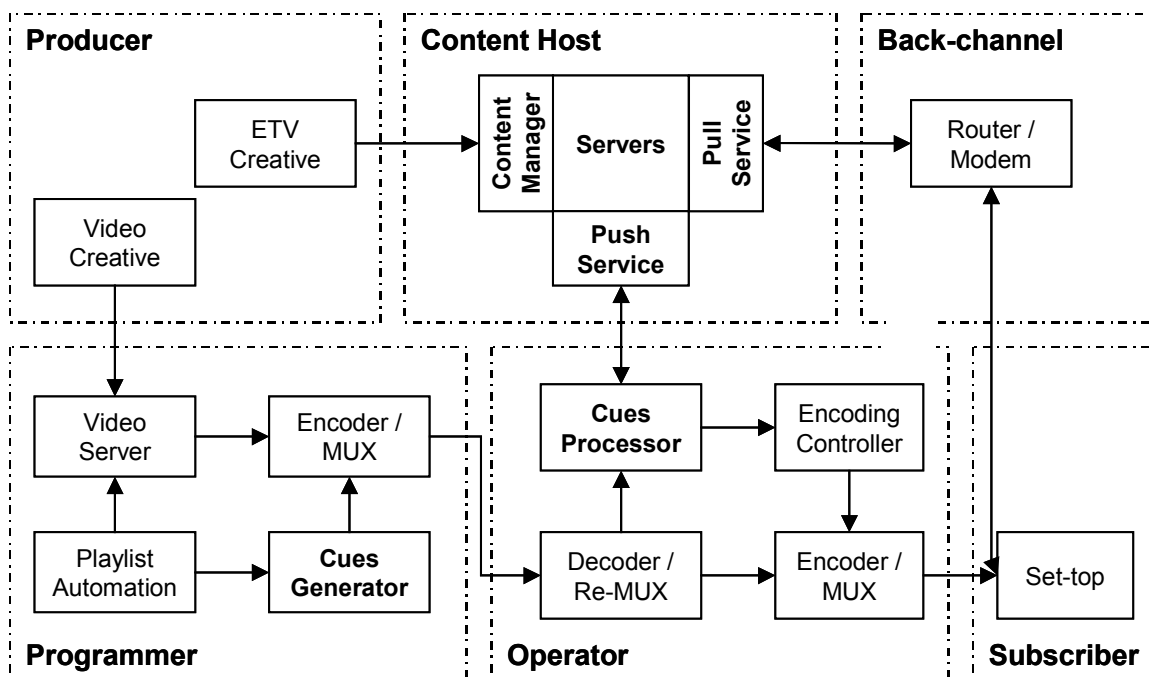
Second, distribution plants have varying requirements and capabilities in forward- and back-channel bandwidths and latencies. A basic choice is whether to embed content in the broadcast signal or send a trigger to load the content over a back-channel; the optimum choice varies by operator and over time even for a single operator.

To address this, we have developed a *Transport Control Architecture (TCA)* where the operators play an active role in managing how the broadcast signal directs content to subscriber equipment. Under the TCA programmers do not send instructions directly to set-tops but rather they send information, or *cues*, that operators use to control their set-tops.

This paper summarizes the TCA and then focuses on its key interface, the protocol for programmers to communicate ETV cues to operators.

TRANSPORT CONTROL ARCHITECTURE

In this paper, we use the term *programmer* to indicate any party that schedules video assets into a channel, and the term *operator* for any party that distributes such channels to consumer equipment. The architecture generalizes to situations where programmers and/or operators are one and the same, or multiple cascading entities such as station groups or meta head-ends.



The figure shows a simplified broadcast chain and back-channel, and the entities that constitute them. Note that the ETV Creative goes not to the programmer but to a *content host* who serves as a clearinghouse for multi-platform content and a server both for operators to load and cache content to be pushed over broadcast, and for set-top back-channels to pull personalized and transmit transaction data over back-channels.

The key components in the figure are the *cues generator* at the programmer and the *cues processor* at the operator. Cue generation takes information from play-list automation and encodes it for use downstream by operators. Cue processing reads such cues and acts upon them, usually by loading and caching content for data carouseling or transmitting a trigger.

The figure shows that the transport of cues is through an embedding in the broadcast signal. Though we anticipate that this will usually be the most convenient means, as we discuss later in this paper, other means are possible and have some advantages.

CUE SYNTAX

Cues may be transmitted using the following general format:

```
!!<URI>{ [attr1:value1] [attr2:value2]...
[attrN:valueN] } { [checksum] }
```

The prefix “!!” indicates the following characters represent a cue. The receiver may choose to act on the cue or to discard it. The remainder of the cue is designed to follow the syntax of triggers from EIA-608B.

The URI identifies the enhancement resource. The URI is enclosed in angled brackets, followed by zero or more pairs of attributes and values in square brackets, then followed by an optional checksum in square brackets. The order of attributes is not important.

As in EIA-608B, the characters included in the cues are interpreted as ISO-8859-1 (Latin-1) characters. Only those within the range 20h to 7Eh are used, the remaining characters should be discarded when processing cues.

Some URIs, attributes or values may require characters not within the above range, or that are excluded in this specification (such as square brackets). These characters should be encoded using the standard mechanism of “%” followed by the two-digit hexadecimal value in ISO-8859-1.

URI Field

The URI identifies a generic enhancement resource that is to be acted upon. An enhancement resource may be a single item such as a trigger, a content package, carousel file or schedule; or it may be a group of several such components.

In this context, a schedule refers to a set of enhancements and their timings throughout a television show or spot. A single cue is able to start the schedule, and the operator takes the responsibility of sending the enhancement components for the show as dictated in the schedule.

The URI can be specified as relative or absolute. Relative URIs allow operators to choose from where to obtain the enhancement data. Allowing the operator to choose the hosting service is a more flexible system, and can help with such issues as load balancing. An enhancement service provider may provide a web interface for obtaining the data; the details of this mechanism are described later in this document.

Attributes

Nine attribute fields are defined in this standard: *component*, *action*, *duration*, *offset*, *id*, *expires*, *source*, *utc*, and *response*.

The *component* attribute is used when the URI references a resource group; the value specifies a component within the group. The URI itself could contain component name, for example, the following cues identify the same resource:

```
!!<adsponsor/enh01>[component:content01]
!!<adsponsor/enh01/content01>
```

Keeping the URI as the root of the group (in this case, *adsponsor/enh01*) enables operators to group cues that apply to the same enhancement. This may be useful when applying actions to the whole group, such as *stop*.

The *action* attribute specifies how the resource is to be processed. The following table defines the current list of actions:

Action value	Action to take
start	Start the enhancement
stop	Stop the enhancement. If the URI specifies a group, stop all its components.
pause	Pause the enhancement and freeze the duration countdown. If the URI specifies a group, pause all its components
resume	Resume the enhancement after a pause.
load	Pre-load the enhancement
cancel	Cancel all previous matching cues. See “id” attribute for details.
query	2-way protocol: Ask for the status of an enhancement
response	2-way protocol: Reply to a status request

If an action is not specified, the default action is *start*.

The *duration* attribute is used with the *start* action. It specifies how long in milliseconds the enhancement should run for. A value of zero indicates the enhancement should run once only. If no duration is specified, the enhancement is run until a *stop* action is called.

The *offset* attribute specifies when the action should be executed. If a *utc* value is specified, the offset is relative to this time, otherwise the offset is relative to the time when the cue was received. Units are milliseconds. If no offset is specified, an offset of zero is used by default.

The *id* attribute, together with the URI and component attribute should uniquely

identify a cue. If a cue is sent multiple times, and these fields remain the same, the previous cues are discarded. For example, a cue with an offset of 30 seconds is sent every 10 seconds, with the offset value decreasing accordingly. The *id* will be the same each time, so the cue is only processed once. Multiple cues sent this way enable systems to cope with occasional delivery failures.

When a *cancel* action is processed, the results depend on whether the three identification fields are specified. If just the URI is specified, all cues with the same URI are cancelled. If the URI and component are specified, all cues matching both fields are cancelled. If the *id* is specified, only those cues matching all three fields are cancelled.

The *expires* attribute specifies in UTC when the cue is no longer valid. Expired cues should be discarded.

The *source* attribute specifies who sent the cue. The will usually be the programmer, but may also be the operator replying to a query in the 2-way protocol.

The *utc* attribute is a UTC timestamp for the cue. The offset attribute is relative to this value.

The *response* attribute is part of the 2-way protocol and contains the response message requested by a query action. If this attribute is present, the action attribute must be *response* and so is ignored.

Checksum

The checksum may be included to detect data corruption. It is identical to the Internet Protocol checksum described in IETF RFC 791 and IETF RFC 1071, also used in EIA-608B. The “!” is included in the calculation. The result is a 16-bit value transmitted as four hexadecimal digits, with the most significant byte first. Characters outside the 20h to 7Eh range are not included in the checksum.

Attribute/value abbreviations

Cue transmission mechanisms often have bandwidth constraints that must be adhered to. The following policy is recommended to minimize bandwidth:

- Remove any attributes that would be present by default. This includes [action:start] and [offset:0]
- Use abbreviations for attributes and action values as shown in the table:

Attribute name	Abbreviation
component:xxx	c:xxx
action:start	a:s
action:stop	a:o
action:pause	a:p
action:resume	a:r
action:load	a:l
action:cancel	a:c
action:query	a:q
action:response	a:e
duration:mmmm	d:mmmm
offset:mmmm	o:mmmm
id:xxx	i:xxx
expires:ttt	e:ttt
source:xxx	s:xxx
utc:ttt	u:ttt
response:xxx	r:xxx

- Keep the URI and component attribute values to a minimum
- Use short id values
- Use short source values

TWO-WAY PROTOCOL

The protocol does not specify a delivery method or platform for the cues. Methods may include delivery in the T-2 service of a video broadcast, conforming to EIA-608 (Recommended Practice for Line 21 Data Services); or out of band delivery via TCP/IP or UDP/IP. The cues do not interfere with, and may be used in conjunction with triggers as in EIA-608B.

The one-way protocol is designed for in-band delivery of cues. It is assumed that in-band cues are synchronized to the broadcast, so that time offsets will be relative to the point in the broadcast when the cue is received.

An out-of-band delivery system may make use of the 2-way version of the protocol. This is an extension to the standard protocol adding the two actions *query* and *response*.

A programmer sends a query action cue asking for the status of an enhancement. The enhancement is identified by the URI, and optionally the component attribute. The operator replies with a response action cue, containing the status of the enhancement in the response attribute. The format of the response value is not defined by this standard and should be defined by the both the programmer and operator involved.

A query action cue containing an id attribute is a request for the status of a cue. The id, URI and component are used to match the cue. The operator replies with a response action cue containing the status of the cue in the response attribute. Again the format of the response value is not defined in this standard.

As mentioned before, a UTC time stamp can be included in cues that are not synchronized to the broadcast. If possible, this timestamp should take into account the time difference between the video feed at the programmer, and the video feed received by an operator. Alternatively, an operator could use the source field to identify the programmer, and adjust the timestamp accordingly.

SCHEDULE RESOURCE FORMAT

A schedule resource specifies a set of enhancement resources with a schedule of when each resource should be processed. The timings will be relative to a point in time, which is usually the beginning of a television show or spot.

A schedule should be specified as a list of cues. All the standard attributes apply except *source* and *utc*. The original cue that referenced the schedule resource determines both these attributes. Also, its offset value will be added to the offset value in each schedule cue.

CUE EXAMPLES

The following are all examples of valid cues:

```
!!<http://itv.adsponsor.com/spi/enh01>[component:content03][action:start][duration:30000]
```

```
!!<adsponsor/enh02>[source:aprogrammer][expires:20010601][offset:10000][id:0077]
```

```
!!<adsponsor/enh03/trigger01>
```

```
!!<adsponsor/enh04>[c:content02][a:l][s:aprogrammer][u:20010101T103922]
```

```
!!<adsponsor/enh04>[c:content02][a:q][s:aprogrammer]
```

```
!!<adsponsor/enh04>[c:content02][r:active][s:anoperator]
```

CUE PROCESSING

This section suggests some procedures that an operator should follow when processing cues.

A local cache is recommended for storing enhancements before insertion. The *load* action indicates enhancements are to be prepared in advance.

A content host's Service Provider Interface (SPI) is a convenient mechanism for supplying enhancement resources. The operator makes a request to the content host specifying the URI, component, and platform. The content host searches for the resource that matches the platform, and returns the required data. A platform may have several possible data formats, so the SPI should also return the specific data type. The cue only specifies the resource identifier; no assumption can be made about the generic data type (such as trigger or content) until the data type is returned by the SPI.

To summarize the process: the cue identifies the generic resource; the operator specifies the platform and requests the resource from a SPI; the SPI returns the resource and its data type.

Each enhanced television platform requires insertion tools and equipment to insert the data into the video stream. There may be several such systems available for each platform. The operator must retrieve the platform specific data and, if necessary, process the data so that it conforms to the insertion system interface.

A content host has the choice of supplying basic platform specific data, or supplying data that conforms to a specific insertion system. If an insertion system becomes standard for a platform then it makes sense for the service provider to provide data in the correct format. This would allow extra control over insertion parameters.

For standard television transmission, automation systems produce *as-run* logs to record the shows/spots that have aired. In the same way, operators should be able to produce as-run logs for enhancements that have been inserted. Each source will have its own associated log.

An operator and the programmers need to specify an error handling policy for the operator to implement. At the very least this should include steps to deal with cue parsing errors, data retrieval errors, and insertion errors. In the future the cue protocol may be extended to include an attribute to specify error-handling policy.

CONTENT HOSTING GUIDELINES

By providing a standard interface, content host service providers are able to serve content for multiple platforms from a single source. The suggested format of this interface is as follows:

ContentHostURL/EnhancementURI/Component

Where *Component* is optional. For example,

`http://itv.contenthost.com/spi/adsp
sor/enh01/content01`

This SPI format allows the operator to select the content host when the enhancement URI is relative, and also allows the programmer to specify the content host with an absolute enhancement URI.

The following table lists the parameters that should also be included in the request. Note that *platform* is the only required parameter.

Parameter	Value
source	Programmer; same value as cue source attribute
client	Operator
platform	Standardized name of the enhanced television platform
device	Standardized name of the insertion system
Any other platform- or insertion-specific parameters	

The content host should attempt to return data that matches the device. If the device is not specified, or is not recognized, the platform should be used to locate the data. If the platform is not supported or recognized, the request should be rejected.

The content host is also able to choose resources based on the programmer and operator. This may be useful for regional targeting of enhancements.

If an enhancement is not available for a platform, the content host has the option of returning a default enhancement.

CONCLUSIONS

This paper has presented a new way for programmers and operators to synchronize in the deployment of Enhanced Television applications. We believe that it is a big win for the operators as it gives them the means to optimize for the network and their choices of delivery platforms. We also believe that it is a big win for the programmers who can achieve maximum ETV distribution with a minimum of distribution-dependent signal-processing.

While we have designed this protocol with some accounting for deployability, we recognize that there are several major open issues, including the following. First, as we mentioned the transport of the one- and two-way protocols between programmers and operators may depend on security and provisioning restrictions, as well as video formats. Second, that operators may need to make restrictions on the use of cues, e.g., that cues must appear at least some offset (as much as a day) before the interactive session.

Third, that this protocol must be more completely reconciled with existing and ongoing standard efforts. And finally, that resolution of the above issues will require bilateral and multilateral agreements among programmers and operators, and at some level this will require an accounting of larger business interests.

REFERENCES

EIA/CEA-608B, Line 21 Data Services (October 2000)

IETF RFC 1071, Computing the Internet Checksum (September 1988)

IETF RFC 2396, Uniform Resource Identifiers (URI): Generic Syntax (August 1998)

ISO-8859-1: 1987, Information processing – 8-bit single-byte coded graphic character sets – Part 1: Latin alphabet No. 1

Systems of Time,
<http://tycho.usno.navy.mil/systime.html>

CONTACT

Dr. Jay Weber
CTO
RespondTV, Inc.
1235 Howard St, San Francisco, CA 94103
jay@respondtv.com

APPLYING NEW TECHNIQUES TO IMPROVE ANALOG-TO-DIGITAL LINK PERFORMANCE—THE ROLE OF DIGITAL SIGNAL PROCESSING

Venkatesh Mutalik, Marcel Schemmann, Craig Malone,
Zhijian Sun, Long Zou, Mary Margaret Ong, Alexei Chkourko
Philips Broadband Networks

Abstract

Analog-to-digital conversion techniques are finding increasing application in HFC transmission networks as a means of transporting analog signals reliably. Among the advantages of these techniques are a range of multiplexing options, mass-produced devices, and the ability to perform additional signal processing in the digital domain.

The role of Digital Signal Processing (DSP) in an analog-to-digital converted CATV system may be divided into two broad sections. In the first, the processing is independent of the specific nature of payload (e.g., the data could be of QPSK, QAM16, QAM64 or QAM256 types). Here, we would primarily be designing general compression schemes as well as automating various self-test and EMS routines, activities which are roughly independent of the digitized information content. In the second, the signal processing can take advantage of known characteristics of the payload data. For example, here, we would be interested in designing appropriate filters, demodulators, and other devices more appropriate for distributed CMTS-type applications.

The authors have implemented a hybrid combination of Digital Signal Processing (DSP) and traditional RF technology to improve HFC return path performance. This approach uses an A/D optical return path transmitter and receiver which, in addition to normal "brute force" coding, adds DSP algorithms that enhance the achievable dynamic range and conserve

bit rate. Test results on prototypes show dramatic improvement in Noise Power Ratio (NPR), and simulations closely match the measured performance.

This paper will present some of these results and address digital signal processing of the payload agnostic type. It will also describe some of the technical details in relation to currently available digital return systems.

INTRODUCTION

In the HFC CATV return band, systems suffer from several noise sources in the cable plant and in the optical links. Noise from the cable plant can be bursty and concentrated at several frequencies [1]. This noise is determined by the existing cable plant and can not be reduced by a transmitter.

However, a transmitter needs modulation headroom before the clipping point to accommodate noise bursts in some channels without affecting other channels. Noise in the optical links stems from laser RIN, shot noise, and receiver thermal noise, and it can be incurred several times if the signal is re-transmitted. State-of-the-art analog transmitters use selected 1310 nm DFB laser transmitters and can achieve an acceptable signal-to-noise ratio with headroom over a limited link length of about 30 km from the node to the hub.

Analog-to-digital conversion followed by digital transmission of the converted data has been proposed as a method to overcome the analog noise sources and thus relax laser requirements, increase transmission distance, and provide a more robust re-transmission capability. Such schemes require very high data rates to achieve an acceptable transmission quality, and results obtained so far at lower rates yield marginal performance.

In this paper, the requirements for the transmission data rate and clock jitter are discussed for a dual RF channel input digital transmitter. Then a DSP algorithm is proposed that allows reduction of the total data rate to fit existing OC-48 standards with excellent transmission quality. The limit for acceptable performance is explored, and results are compared with test data. We further discuss the link budget requirements and DWDM system requirements for a fiber-conserving system that operates on the ITU grid. The data reduction provides enough room for the addition of bits for FEC and DC balancing, which results in an extremely high link budget and transmission distance. Digital control provides excellent wavelength stability up to and exceeding an 85° C module temperature environment, with reduced power dissipation.

System description

The return transmitter operates from a node. In the field, internal node temperatures in the range of -40° C to +85° C can occur. In a system where, for example, 32 nodes are to be combined into a single fiber with 0.8 nm (100 GHz) spaced ITU channel lasers, a ± 0.1 nm wavelength stability is required for the transmitters.

Use of a DWDM combiner in the field would put very strict requirements on its temperature performance. For this reason,

we propose using a passive coupler. For 32 channels, a 12 dB loss is to be anticipated. The maximum fiber span for modem operation is assumed to be 100 km, which yields a 20 dB fiber loss (0.2 dB/km). Finally, in the headend, a DWDD is assumed with a 6 dB implementation loss. This yields a total loss of 38 dB. Figure 1 depicts the described system. A comparison of the two scenarios [2, 3] indicates savings on transmitter receiver pairs, wavelengths, capacity, cost of the DWDM, and the cost of erecting and maintaining the hub.

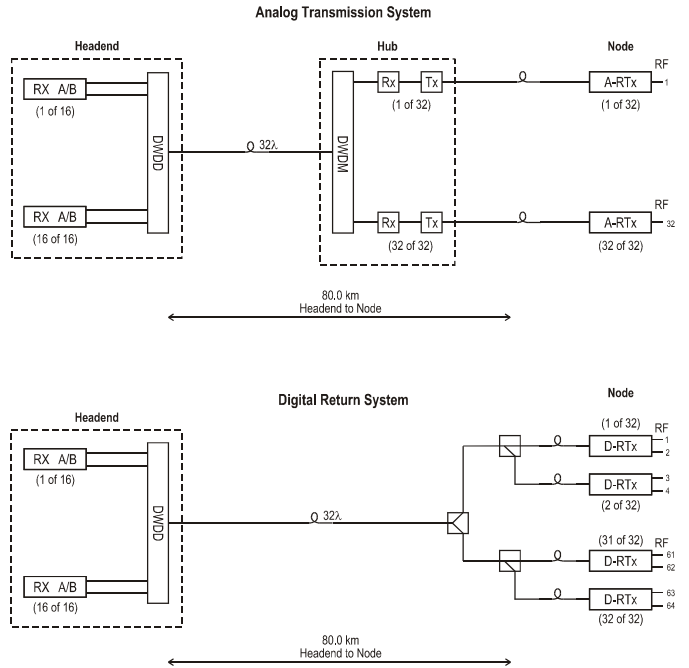


Figure 1. Comparison of analog and digital return transmission system architectures

Basic implementation

Optically, this link requires a low chirp DWDM OC-48 laser with a 7 dBm average output power and typically 1 dB dispersion penalty at 100 km. The receiver requires an APD detector with a -32 dBm sensitivity to obtain the required budget and link length.

The analog to digital converters digitize two RF input channels operating in the 5.42 MHz band (or 5.48 MHz in FTLA systems). The minimum sampling rate is twice the RF bandwidth and thus should be above 100 MHz. A typical requirement on the attainable signal-to-noise ratio as determined by an NPR test is 40 dB. This NPR is to be maintained in a dynamic window of at least 15 dB [4]. This is investigated in the next section.

THEORY AND SIMULATION

Effective Modulation Depth

Optical, RF, and digital performance are all related by a parameter called the Effective Modulation Depth (μ_{eff}) for the A/D converter. This μ_{eff} is a counterpart of the more familiar Optical Modulation Depth (OMD) frequently encountered in optics literature. Just as the laser is a compressing entity in the optical domain, so too is the A/D converter to the digital system. It may be noted that for digital optical transmission, an OMD exists as well for the laser, but a well-designed transmission system does not have any impairment due to the OMD of the laser as a function of input RF power.

In order for the system to obtain some required performance, a minimum Effective Number of Bits (ENOB) is required from the A/D converter. Note that this number is smaller than the Number of Bits (NOB) of the A/D converter due to non-linearities in the A/D converter [5].

As illustrated in Figure 2, for a multi-channel RF input signal, an effective modulation depth (μ_{eff}) of close to 30% leads to significant clipping distortion.

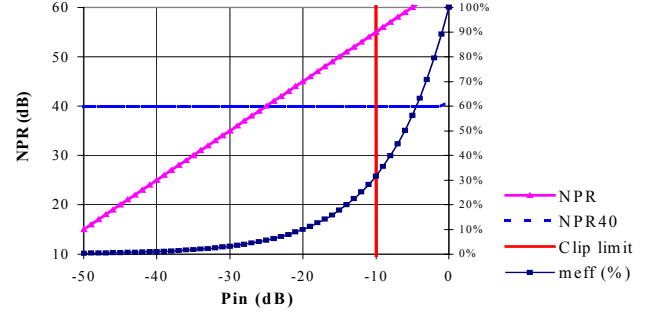


Figure 2. NPR and μ_{eff} as a function of RF power

To meet the 40/15 NPR requirement, one should try to obtain a 40 dB NPR at 15 dB below the A/D converter's clipping point. Hence, for $\mu_{\text{eff}} = 30\%$, a minimum signal-to-noise-ratio of 55 dB is required (disregarding clipping impairments). This translates to 62 dB for a single tone with a 100% amplitude modulation index (71% μ_{eff}) and requires an ENOB of 10. Consequently, the NOB must be 11 or higher.

The sample and hold (S/H) timing jitter should be approximately 15 ps to prevent ENOB degradation from sampling jitter. In order to satisfy the ENOB requirements, a 12-bit AD converter is needed. A sampling rate of 103.68 MHz is used to simultaneously meet the Nyquist condition and be a perfect match for the OC48 transmission rate. The total data rate is then calculated as follows:

$$2 \text{ channels} \times 12 \text{ bits} \times 103.68 \text{ MHz} = 2.488 \text{ Gb/s}$$

This is the exact OC48 standard data rate. Hence, cost-effective components are available for the link. Notice, however, that there is no overhead left for word synchronization, DC balancing, error correction, and status monitoring.

Simulation routine

The system performance of a system with multiple RF carriers and with an NOB

of 12 and a sampling jitter of 15 ps rms can be numerically estimated by sampling and digitizing a set of carriers with 2^{ENOB} levels. By taking the Fourier transform of the resulting set of data, a spectrum results.

In the input spectrum, a slot is left out. The timing jitter and the quantization noise generate spectral content in the slot. Comparison of that power to the carrier power yields an estimate of the anticipated NPR and is displayed as the output of the simulation.

DIGITAL SIGNAL PROCESSING

Thus far, a general description of the digitizing and recovery process has been explained. It is well known, however, that some overhead is required in digital transmission to effectively transmit information. The only possible way out of such activity is by somehow reducing the amount of required information payload from the current OC48 data rates. Doing so, however, requires either a reduction in the number of bits or a change in the sampling frequency, neither of which is conducive to required performance. Therefore, some techniques have to be applied to protect the performance as well as to allow some leeway for data processing.

Test results

The number of bits required to digitize the input signal can be reduced by suitably processing the signal before it is digitized such that a converter (e.g., 10-bit AD) can be applied. However, this leads to large distortions, so an inverse operation is required at the receiver side.

The input signal range is reduced at the transmitter side. The operation is inverted at the receiver side. This results in

an enhancement of the quantization noise for large input power such that degradation in SNR is expected. However, for large input power, the signal is very far above the quantization noise, so the performance is acceptable.

A serious lingering problem with this method is that to maintain a 40 dB NPR, an extremely good match between the analog compressing and de-compressing circuitry is required, which is not practical over the operating temperature range. For this reason, 12-bit A/D and D/A converters have been chosen, and the DSP algorithm operations are done in the digital domain. Using the method outlined above, expected NPR performance is calculated for compression from 12 to 10 and 8 bits.

Obviously, a 10-bit algorithm scheme is possible with significant margin over the required NPR performance. Even down to 8 bits, the required NPR performance can be met. The efficient algorithm scheme frees up bits that can now be used for overhead. These are now used for addition of FEC and DC balance of the code. NPR test results for 12 to 10 and 12 to 8 bit compression are shown in Figure 3 for a transmitter with 7 dBm output power (10 dBm peak) and an APD receiver on a 100 km fiber link.

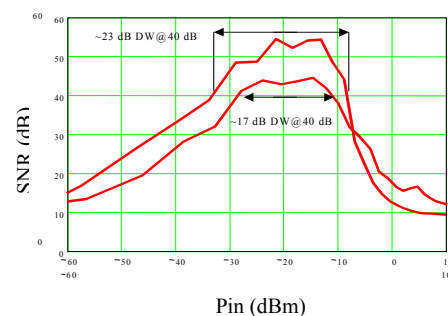


Figure 3. Estimated NPR performance of a digitally processed system

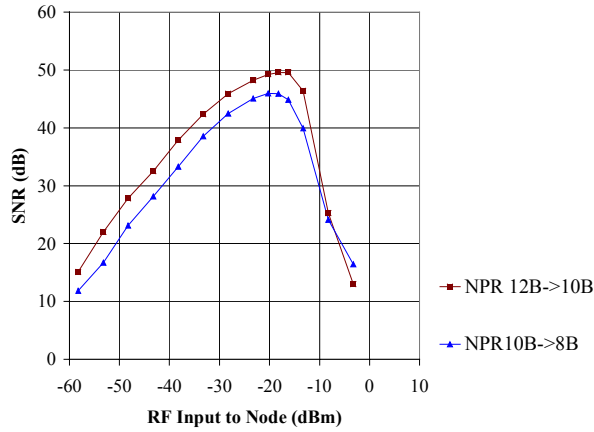


Figure 4. NPR test results for 12 to 10 and 10 to 8 bit digitally processed systems

Figures 3 and 4 reveal excellent agreement between model and test data; for a 12 to 10 bit transformation, 40 dB NPR can be obtained in a 20 dB dynamic window, which well exceeds the requirement.

The curve presented in Figure 5 shows performance at different link budgets using FEC for the digital link. Using the available overhead and signal processing allows operation well above the required link budget. No performance degradation is observed up to 41 dB budget, and the link is maintained up to 44 dB budget. The FEC allows monitoring of the error rate, and the system is switched off if a 20 dB SNR cannot be maintained at the clipping side.

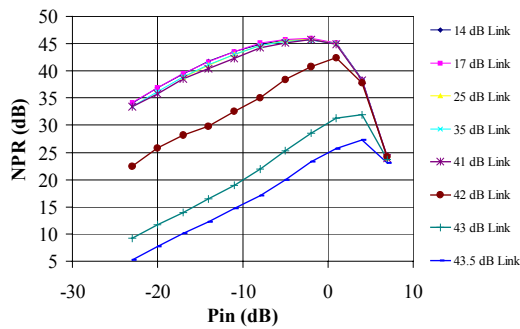


Figure 5. NPR test results of a 12 to 10 bit digitally processed system over various link budgets

In Figure 6, QPSK and 16 QAM error rates are plotted for comparison with the NPR curves [6]. On the noise side, the BER and SNR curves yield the same result. On the clipping side, the BER is degraded at a lower power level due to the non-Gaussian nature of clipping noise.

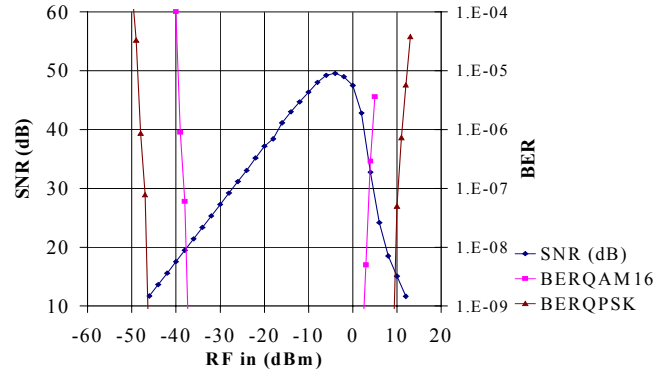


Figure 6. NPR and BER test results of a 12 to 10 bit digitally processed system

Digital control of laser temperature

The presence of a powerful DSP processor also allows intelligent control of the laser Peltier element in a full bridge switch mode power supply configuration. The DSP can be programmed to minimize dissipation and switching noise through intelligent control of the switches. The resulting wavelength stability and power consumption is plotted in Figure 7.

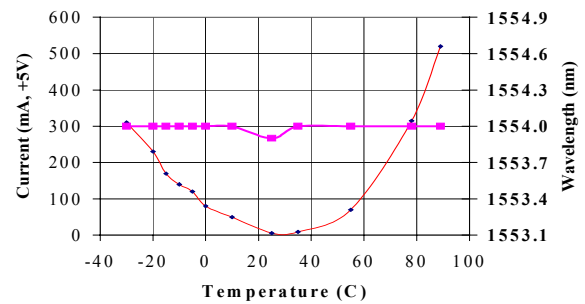


Figure 7. Test data of wavelength and power consumption as a function of module temperature

Clearly, the laser wavelength can be maintained well within limits up to 90° C module temperature. The power consumption that is required to cool the laser is as low as 3W up to 85° C module temperature.

The NPR performance of the system over temperature shown in Figure 8 demonstrates the power of DSP in maintaining system performance over temperature. An ambient node housing temperature of 60° C presents itself as 85° C for the transmitter module.

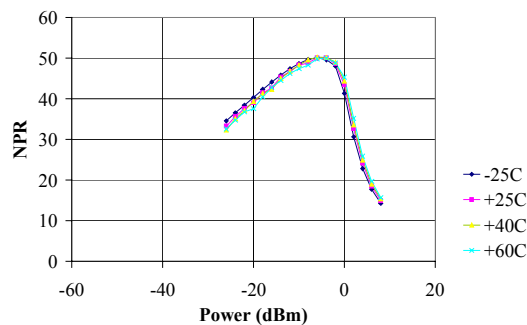


Figure 8. NPR test results of a 12 to 10 bit digitally processed system over ambient node temperature

The DSP monitors the laser temperature and switches off the laser power in case the wavelength is out of range; thus, neighboring channels will not be affected. It also monitors the cooler current consumption and limits that to prevent overloading the power supply.

CONCLUSION

A dual channel AD CATV return band transmitter was presented operating on the ITU grid. Requirements were derived for an EDFA free operation of a return system that combines 64 individual return channels on a single fiber and transports these over 100 km of single mode fiber.

It was shown that at least a 12-bit A/D converter is required for acceptable operation, resulting in an excessive requirement of the total data rate. An information efficient algorithm was then implemented with 10-bit transmission in a DSP to obtain 40 dB NPR in a 20 dB dynamic window combined with code balancing.

An additional algorithm is proposed that allows a 40 dB NPR in a 15 dB dynamic window with only 8-bit transmission. The digital link FEC allows a link budget in excess of 40 dB and link lengths in excess of 150 km. Further application of the DSP to the control of the laser cooler has yielded better than 0.1 nm wavelength stability in the -40° C to +90° C degree temperature range at a power consumption of only 3 W at 85° C module temperature.

Similar success with respect to NPR performance has been achieved with un-cooled DFB lasers with distances approaching 80 km.

Clearly, the application of DSP can significantly improve the performance of digital return systems

ACKNOWLEDGEMENTS

It is a pleasure to acknowledge Sudhesh Mysore and Rick Villa for their help in determining the NPR characteristics of this system at the AT&T Broadband Labs.

REFERENCES

1. K. D. LaViolette, T. Staniec, "Lessons for the Interactive Return System," NCTA Technical Papers, 1996, p. 227.
2. O. Sniezko, et al., "HFC Architecture in the Making—Future Proofing the Network," NCTA Technical Papers, 1999, p. 20.
3. S. Mysore, R. Villa, G. Beveridge, V. Mutalik, "Choosing the Right DWDM Design—How to Evaluate Variables Affecting Bandwidth and Cost," CED, July 2000, p. 88.
4. O. Sniezko, T. Werner, "Return Path Active Components Test Methods and Performance Comparison," SCTE Conference on Emerging Technologies, Jan 1997.
5. B. Sklar, Digital Communications: Fundamentals and Applications, Prentice Hall Inc., 1998.
6. R. Thomas, J. Monroe, "Reverse Path Characterization: A Summary of Field & Lab Test Results," SCTE Proceedings Manual, 2000, pp. 475-491.

BNN -- A COMPREHENSIVE BANDWIDTH MANAGEMENT TECHNIQUE FOR THE FORWARD PATH

Jerry Monroe, Philips Broadband Networks
Ray Thomas, Time Warner Cable

Abstract

Standardization of digital capabilities of existing and to-be deployed HFC networks calls for advanced techniques of measuring digital performance.

The authors quantify high-speed forward (downstream) data capabilities through the use of the Bit Error Rate in the Noise Notch (BNN) technique.

Additive White Gaussian Noise (AWGN) is used to simulate full forward path digital loading on a selected portion of a network under test. Inserted into this fully loaded bandwidth is a non-forward-error-corrected 6 MHz, 256 QAM evaluation signal.

Signal levels are varied, and as a result, the true bit-by-bit dynamic range of operation is fully understood. Additionally, the effects of the digital tier on analog channels can also be measured by this technique.

Benefits for the manufacturer are a repeatable test, which quantifies predicted performance and can easily be demonstrated for network operators. Benefits for the operator include the ability to certify and to re verify performance of recently deployed networks.

This technique can also be used successfully in predicting performance of legacy networks when a fully loaded digital tier is to be added.

The authors present lab tests using this technique with analog loading provided by CW carriers or analog headend signals.

By using this technique, bandwidth management of the forward path can be easily achieved, since the dynamics of the path will be understood. This tool will be very important to deployed networks and their need to know how much headroom is available for future digital services.

INTRODUCTION

Rigorous qualification testing of the electronic components in a broadband network was not a consideration in the early days of cable. The immediate goal was to fabricate a device that did the job and get it into service as soon as possible. As our industry has evolved, the need for qualification testing has been discovered and refined. Today, measurement of characteristics like Mean Time To Failure (MTTF) and Mean Time Between Failure (MTBF) are standardized and routine.

Through testing, manufacturers seek to demonstrate to network operators that their products are superior to those of other vendors, and operators seek hard proof that the equipment they are buying meets high standards for performance, quality, and reliability.

This paper shares the results of on-going work to develop test and qualification methods that accurately simulate and predict the performance of RF amplifiers, optical transmitters, and optical receivers in broadband networks transporting analog and digital signals to deliver advanced two-way services.

TESTING TODAY

What Drives the Way We Test?

Currently, several different types of test signal loading are used to conduct qualification testing, with manufacturers using different test configurations and signal loading depending on the network operator's requirements. In nearly all situations, the test requirements are specified by the network operator.

Broadband network operators are driven to specify various qualification testing methods to meet their system design needs. These needs vary widely, and may include factors such as system size, available bandwidth, channel type, channel loading, dollars invested per home passed, business philosophy, and engineering management philosophy. The goal of qualification testing is to identify products that meet performance, reliability, and cost requirements.

How We Currently Test

A channel plan widely used in broadband networks for forward transmission places analog channels from 50 to 550 MHz and digital channels from 550 to 750 MHz.

Standard practice for qualification testing of active devices in the forward path fills the spectrum from 50 to 550 MHz with Continuous Wave (CW) analog carriers. However, a number of different methods are in use for simulating the digital channel loading of the spectrum from 550 to 750 MHz.

One method of simulating the digital load consists of using a bank of individual QAM modulators (64 QAM or 256 QAM). Loading the 200 MHz spectrum with 6 MHz

digital channels in this manner requires 33 separate QAM modulators.

A second method of simulating the digital channel loading is to mix the output of a small number of QAM modulators with upconvertors/RF generators to create a wide band of digital loading.

A New Way to Test

This paper presents a third method to simulate digital loading for forward network testing: use Additive White Gaussian Noise (AWGN) to fill the spectrum from 550 to 750 MHz.

The results of testing using a combination of analog carriers and AWGN are shown in this paper. Our tests examined AWGN levels representing full digital channel loading that varied with respect to the CW carriers. The various levels of AWGN loading were selected to represent typical operating levels used in broadband networks throughout the world.

Our tests also simulate the changes that will take place as the high end of the spectrum carrying analog signals is converted to carry digital channels. We do this by expanding the bandwidth loaded with AWGN signals to 500 - 750 MHz, while at the same time reducing the bandwidth loaded with analog channels to 50 - 500 MHz. Each digital channel is capable of carrying the equivalent of up to 10 analog NTSC television channels.

The tests documented here seek to determine the performance of forward path DFB optical transmitters when fully loaded from 50 to 750 MHz with analog carriers and AWGN. BER performance was measured as a function of RF drive level to a transmitter under a variety of analog to digital schemes.

USING BNN FOR FORWARD TESTING

The Bit Error Rate in the Noise Notch technique (BNN) is not a new idea. Previously, this technique has been used successfully to quantify bit error rates of common digital modulation modes (QPSK and 16 QAM) used in the reverse path under a fully loaded bandwidth condition. The additional benefit of this type of testing is to validate and demonstrate dynamic ranges of operation of various components of the HFC network. (See R. Thomas and J. Monroe, "Reverse Path Characterization: A Summary of Field and Lab Test Results," SCTE Cable-Tec Expo 2000 Proceedings Manual, 2000, pp. 475-491.)

Overview of the BNN Technique

In general, the BNN technique works as follows. First, define the bandwidth of operation. For the forward path testing, this was chosen as 54 to 750 MHz. Then, fully load the bandwidth of operation with analog and digital signals.

The analog portion of this bandwidth is typically loaded with NTSC-based video channels, spaced at 6 MHz from 55.25 MHz to 547.25 MHz. For lab testing, these channels are simulated with CW carriers generated by a stable crystal-controlled signal generator source.

Complete analysis of forward network performance, requires loading the bandwidth with the full digital component also. To load the 550 to 750 MHz bandwidth we used AWGN because of its large peak-to-average power ratio, which is typical of the 64 and 256 QAM modulation techniques used in today's forward path. This fully loaded condition is shown in Figure 1.

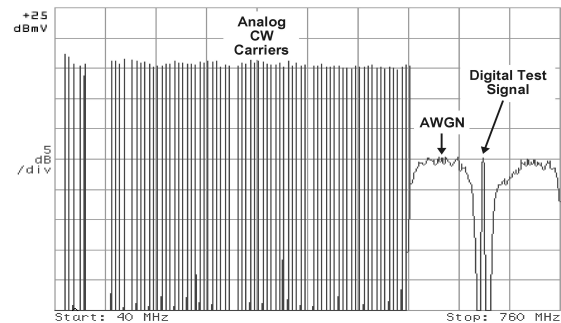


Figure 1.

Fully loaded forward test has analog CW carriers and AWGN to simulate digital loading

Note the "notched" out portion of the AWGN bandwidth. We inserted a 6 MHz wide 64 or 256 QAM test signal into this notch. The IF-based test signal is "up converted" to the desired frequency--in this case 651 MHz. This non-Forward Error Corrected (FEC) test signal is monitored for bit errors. The data signal uses a Pseudo Random Bit Sequence (PRBS) which allows for precise bit-by-bit evaluation. The HFC network, or a portion thereof, can now be evaluated for BER in addition to traditional performance parameters such as carrier-to-noise, composite triple beat and composite second order.

Test Setup

As shown in Figure 2, the test setup for evaluating performance in the forward system calls for three signal sources: a CW signal source; a band-passed, filtered broadband AWGN source; and a 64 or 256 QAM modulator.

Output of the CW signal source is combined with output of the AWGN signal source. The combined signal is notched at the test frequency and has the QAM signal to be evaluated for BER inserted into it. In this test setup the QAM signal is originated from a 43.75 MHz IF-based modulator,

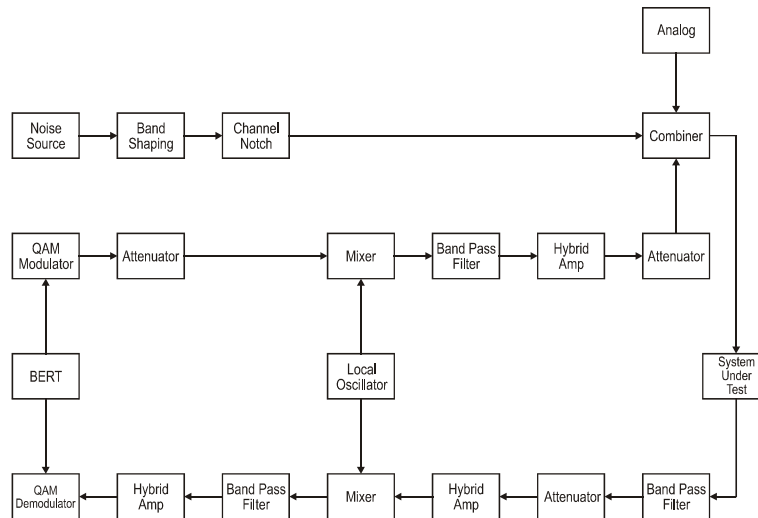


Figure 2.
Setup for using BNN for forward testing

which is unconverted and bandpass filtered to the correct operating frequency. A similar means of down conversion is used to feed a matching demodulator also operating at 43.75 MHz. Both the QAM modulator and demodulator are data and clock driven from a common digital transmission analyzer.

Setting Levels

Setting of digital levels relative to the analog video levels is critical. The following technique is valid for a 50 to 750 MHz amplitude flat (zero dB tilt) test signal only.

First, determine the desired operational levels (For example, digital at -8 dB relative to analog, CW carriers $+15$ dBmV, digital carriers $+7$ dBmV). Next, use the band power marker function of a spectrum analyzer to set the digital band power of the QAM channel accordingly. Then use the same technique to set the remainder of the digital bandwidth. Since the AWGN signal is essentially flat, a representative 6 MHz portion of this bandwidth is all that is necessary for establishing the correct level. See figure 3.

A proof of proper digital level relative to analog level can also be performed as follows.

Analog signal level	=	$+15$ dBmV
Digital signal level	=	$+7$ dBmV
Digital signal bandwidth	=	200 MHz
Number of 6 MHz	=	33
digital signals in		channels
200 MHz bandwidth		$(200 \div 6)$
Theoretical total bandpower	=	22.2 dBmV
of digital bandwidth		
$(10 \log 33 + 7$ dBmV)		
Actual measured bandpower	=	22.5 dBmV
of 550-750 MHz using		
spectrum analyzer		

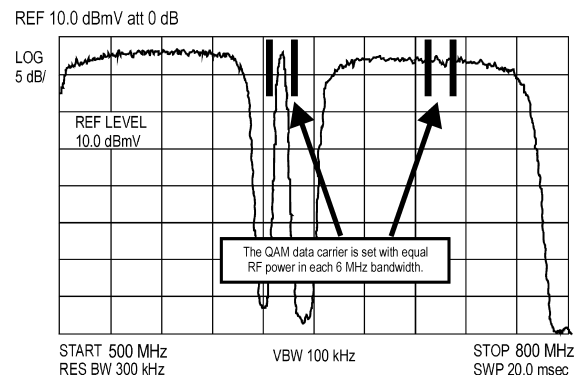


Figure 3.
AWGN with QAM data carrier in the notch

Comparing Apples to Oranges

The power and simplicity of this test is that it allows the user the freedom to use an AWGN source generator to simulate a full complement of digital modulators for a given bandwidth.

However, to verify that AWGN does not present a worst case or less than worst case condition, the following correlation testing took place. A bank of twenty-two 64 QAM generators was used in a direct comparison to an identical bandwidth of AWGN. In each bandwidth scenario, the same 64 or 256 QAM modulated test signal was evaluated, as shown in Figures 4 and 5.

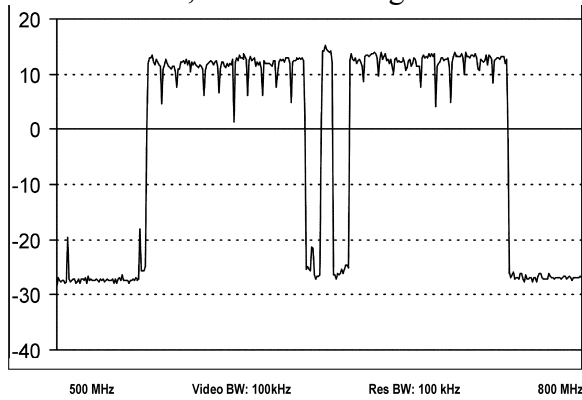


Figure 4. BER measurement of QAM channel + 22 x 64 QAM

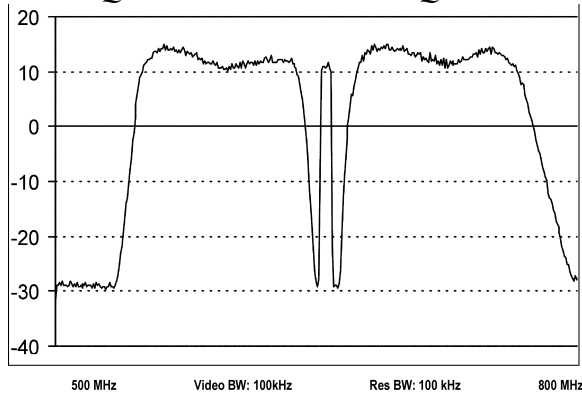


Figure 5. BER measurement of QAM channel + AWGN

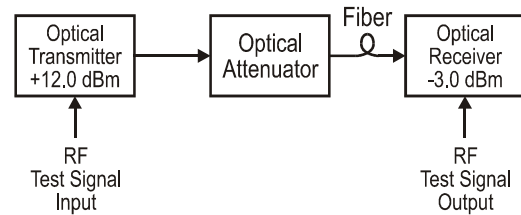
Performing this test on a forward path DFB optical transmitter showed identical BER results when using either power-loading scheme; therefore, AWGN is a valid signal source.

VALIDATING THE BNN TECHNIQUE WITH SAMPLE TESTS

To validate the appropriateness of the BNN technique, we used it for nine sample tests measuring forward path performance. These tests successfully used the BNN technique to characterize BER of a forward path transmitter. This evaluation was chosen because of the increased scrutiny of laser clipping and its impact on digital transmission.

Setup for Sample Tests

As shown in Figure 6, the test setup used a +12 dBm transmitter to feed a 15 dB link composed of an optical attenuator plus 10 km of single mode optical fiber. This link was then evaluated under a variety of RF input conditions.



*Figure 6.
Setup used for tests to
validate BER technique*

Analog Test to Baseline the Transmitter

The transmitter was tested with the 78-channel CW load from 50-550 MHz and the digital load carried at -6 dB.

CTB, CSO and C/N were measured to determine the relationship between the transmitter under test and clipping or the point at which the distortion levels change in a non-linear fashion. Results of this test are shown in Figure 7.

Observations: Clipping of the transmitter takes place at an analog carrier drive level of +18 dBmV. This is indicated by the non-linear increase in both CTB and CSO for a given dB increase in drive level.

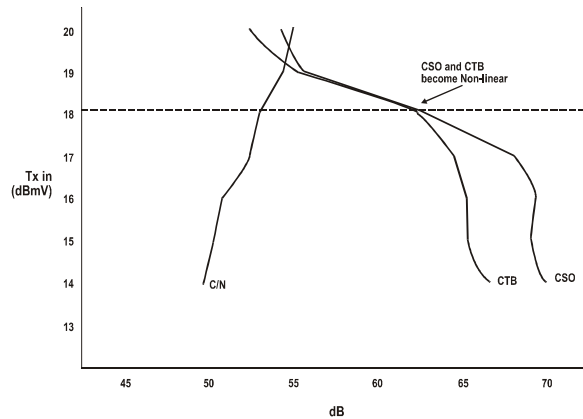


Figure 7.
Transmitter input level vs CTB, CSO, C/N

BER Test with Digital at -6 dB

The test transmitter was evaluated for BER for both 64 and 256 QAM. Transmitter input level of the digital signal remained at -6 dB relative to the analog carriers. Measured BER was recorded as a function of RF drive level. Results are shown in Figure 8.

Observations: Analog techniques indicate that clipping occurs at a drive level of +18 dBmV; however, minimum BER degradation does take place at +15 dBmV for 256 QAM and at +17 dBmV for 64 QAM.

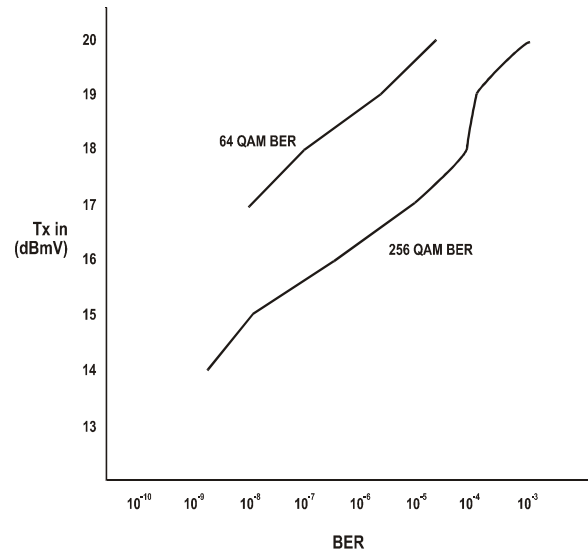


Figure 8.
Transmitter input level vs BER
with digital at -6 dB

BER Test with Digital at – 8 dB

In this test, transmitter BER was evaluated for both 64 and 256 QAM. Transmitter input level of the digital signal was –8 dB relative to the analog carriers. Measured BER was recorded as a function of RF drive level. Results are shown in Figure 9.

Observations: Results are very similar to when digital operated at –6 dB down, with negligible variation in 64 and 256 QAM.

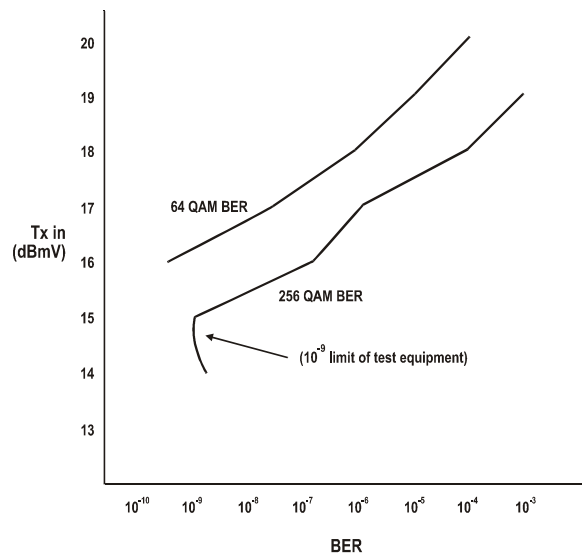


Figure 9.
Transmitter input level vs BER
with digital at – 8 dB

BER Test with Digital at –10 dB

In this test transmitter BER was evaluated for both 64 and 256 QAM. Transmitter input level of the digital signal was –10 dB relative to the analog carriers. Measured BER was recorded as a function of RF drive level. Results are shown in Figure 10.

Observation: Results are again similar to when the digital is operated at –6 and –8 dB down.

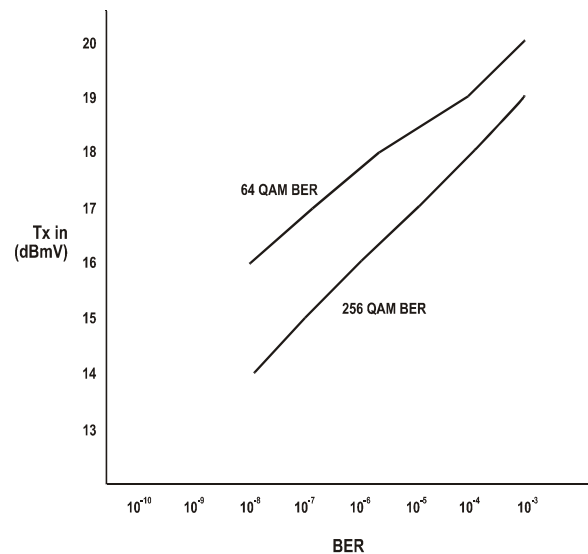
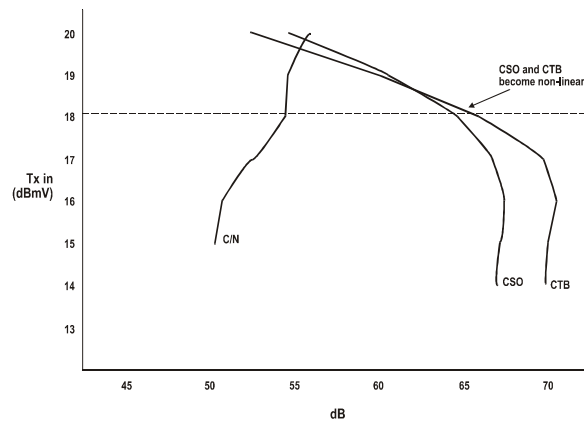


Figure 10.
Transmitter input level vs BER
with digital at –10 dB

Analog Test to Baseline the Transmitter with Expanded Digital at – 6 dB

Input signal was modified to 69 channels of analog (50-500 MHz) and digital loading from 500 to 750 MHz at –6 dB relative to the analog carriers. The transmitter was baselined. Results are shown in Figure 11.

Observation: The clipping point with 69 analog channels remains within a dB of the clipping point measured when loaded with 78 analog channels.

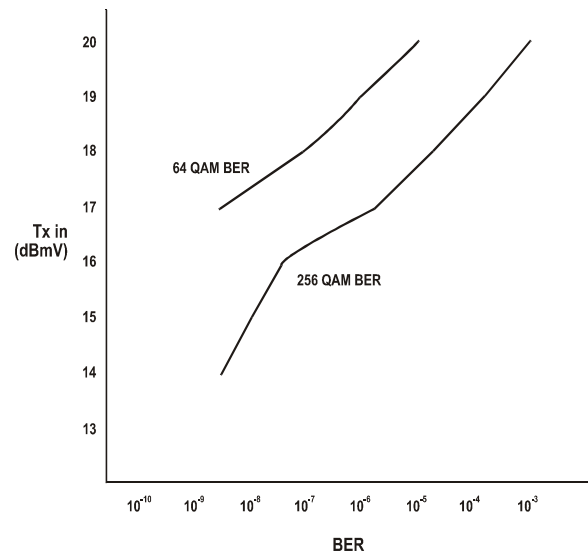


*Figure 11.
Transmitter input level vs CTB, CSO, C/N
with analog 50–500 MHz
and digital 500–750 MHz*

BER Test with Expanded Digital at –6 dB

In this test, transmitter BER was evaluated for both 64 and 256 QAM. Transmitter input level of the digital signal remained at –6 dB relative to the analog carriers. Measured BER is recorded as a function of RF drive level. Results are shown in Figure 12.

Observation: The 64 and 256 QAM performance with 50-500 MHz analog and 500-750 MHz digital was similar to the performance with 50-550 analog and 550-750 digital.



*Figure 12.
Transmitter input level vs BER
with analog 50–500 MHz
and digital 500–750 MHz at – 6 dB*

BER Test with Headend Analog and Digital at – 6 dB

This test used a live, high-quality, direct cable television headend feed to provide analog loading from 50-550 MHz. The digital load was inserted at –6 dB relative to the analog carriers. The system was tested again for 64 and 256 QAM BER. Results are shown in Figure 13.

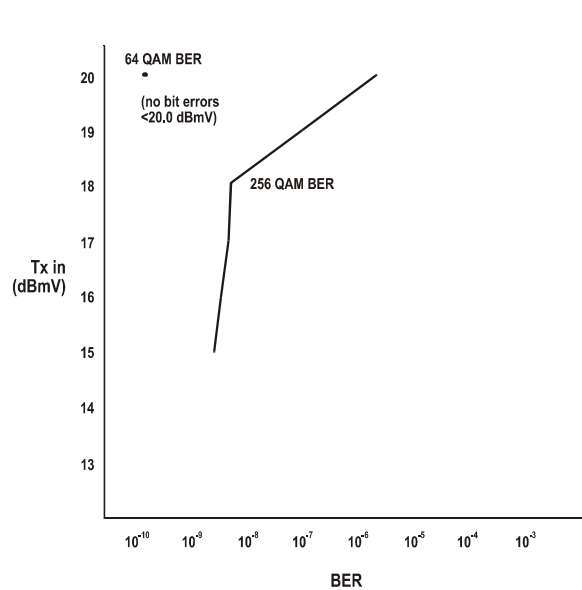


Figure 13.
*Transmitter input level vs BER with headend
analog and digital at – 6 dB*

BER Test with Headend Analog and Digital at – 8 dB

This test used local cable television headend as the analog channel feed for 50-550 MHz and digital load inserted at –8 dB relative to the analog carriers. The system was tested again for 64 and 256 QAM BER. Results are shown in Figure 14.

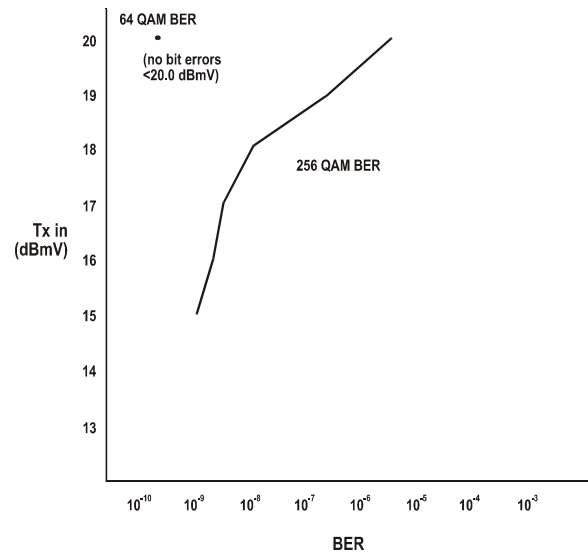


Figure 14.
*Transmitter input level vs BER
with headend analog and digital at – 8 dB*

BER Test with Headend Analog and Digital at – 10 dB

This test used local cable television headend as the analog channel feed for 50-550 MHz and a digital load inserted at –10 dB. The system was tested again for 64 and 256 QAM BER. Results are shown in Figure 15.

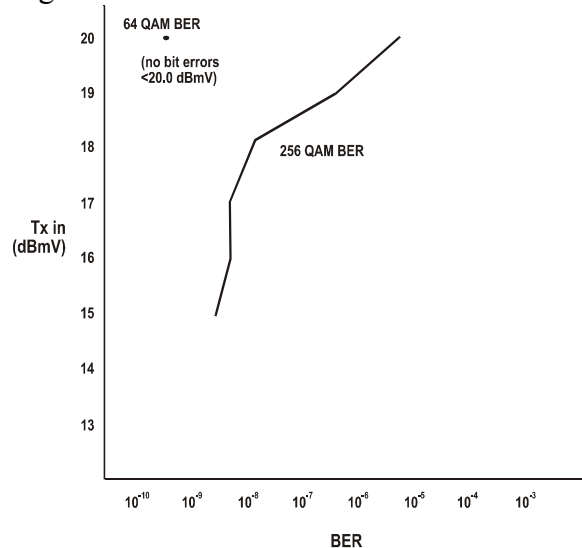


Figure 15.

*Transmitter input level vs BER
with headend analog and digital at – 8 dB*

Observation: All three tests using headend signals as the analog feed show a significant improvement in BER for both 64 QAM and 256 QAM, as compared to an analog feed of CW carriers of the same peak level.

CONCLUSIONS

The actual testing results show the BNN technique to be a valid means for measuring the performance of the forward path in networks carrying a combination of analog and digital signals.

When applied to the forward path, the BNN test technique clearly yields results that very accurately document bit error performance. This technique relies on readily available resources and inexpensive filters for band shaping to easily create and

duplicate a variety of test signal loading conditions.

Manufacturers and system operators can use this technique for system qualification and optimization with confidence in the results derived.

Initial testing indicates that the optical DFB-based transmitter is clipping at drive levels approaching +18 dBmV with the digital load operating at –6 dB relative to analog. Manufacturer specified drive levels for this particular device are at +15 dBmV

It is interesting to note that measurable 256 QAM degradation of BER does appear to take place several dB before measured CTB and CSO clipping (and to a lesser degree 64 QAM).

However, even at the drive level of +17 dBmV, for which the CTB and CSO were observed to be linear, both 64 and 256 QAM are still not at bit error rates of $E10^{-4}$ which is considered a worst case operating level for digital set top boxes and cable modems.

Removing the top eight NTSC channels and increasing the digital tier by another 50 MHz was demonstrated to have minimal effect on 64 and 256 QAM BER, as the actual clip point remains essentially the same.

Most interesting is the final set of tests where the actual CW analog carriers are replaced by high-quality live analog video carriers originated at the headend. A significant decrease in optical modulation index was noted, and as a result associated BER for 256 QAM modulation was reduced.

This test should give confidence to network operators that negligible or no clipping of the laser occurs when a full tier of digital signals is added to the forward path. (Provided, of course, that the transmitter is operated per typical manufacturer recommended drive levels).

BRIDGING AND MANAGING QOS - CABLE AND HOME NETWORKS

Mariano Schain Itay Sherman
Texas Instruments – Cable Broadband Communications

Abstract

In order to enable end-to-end Quality of Service (QoS) for different services, such as voice and video, a mapping of the QoS mechanisms of any two bridged networks must be supported and configured. A direct mapping approach is not scalable or feasible. In this paper we present a generic QoS scheme. We present in detail the QoS framework of DOCSIS1.1 and provide an overview of the QoS support in current and future versions of home networking protocols and architectures (e.g. 802.11, HPNA). For the frameworks presented, we explore the method to map their QoS mechanisms to the generic scheme, thus allowing scalable indirect QoS bridging.

INTRODUCTION

Cable operators are facing a new challenge in distributing their services into the home.

The DOCSIS (Data Over Cable System Interface Specification) standard has been widely adopted by the cable industry as the prevailing protocol for Internet over cable. The first version of the specification (DOCSIS 1.0) mechanisms was targeted at basic data transmission, providing a best-effort level of service. The new version (DOCSIS 1.1) is built on top of the previous 1.0 specification and enables operators to provide consistent and reliable digital services (such as Voice and Video) through the use of sophisticated QoS (Quality of Service) and Network Management.

However, with the emergence of multiple Home Networking technologies, operators find themselves using technologies over which they have significantly less control to distribute their services throughout the home. While these technologies can enhance the customers' experience, they also potentially put the reliability of cable services at risk. To further complicate things, many Home Networking technologies are being offered to consumers: wireless, powerline, powerline and new wire technologies. In most of these categories, several technologies are competing with each other. Furthermore, the list of competing technologies is getting longer as new technologies emerge. It is expected that several Home Networking technologies will find their way into the home, resulting in a heterogeneous home network.

In such a heterogeneous environment, a data flow carrying service information may cross different home network segments on its end-to-end path. The service quality can't be guaranteed unless it is configured for each and every segment. Special consideration is given to the bridging points – the network entities that attach different home network segments together. These network bridges are responsible for forwarding the data flow across their network interfaces so that it seamlessly crosses the segment's boundary without affecting the quality of service.

To achieve the QoS-bridging goal, one may attempt to map the QoS mechanisms of any two protocols 'to be bridged.' This approach, however, is not always possible -

different network protocols provide different levels of QoS. This approach also is not always efficient because too many parameters must be considered in the mapping. Moreover, the scalability of such an approach is more than questionable because, for every new network protocol, a mapping to all existing protocols should be provided.

A generic QoS specification, as described in this paper, is both feasible and scalable as no direct mapping is required. Every network technology is mapped to the generic QoS specification, requiring a linear total number of mapping rather than square.

The paper is organized as follows: first, we present in detail the QoS framework of DOCSIS1.1. Next, an overview of the different home-networking technologies is given, including their strengths and weaknesses and for which applications they are most well-suited. The next section presents the concept of generic QoS specification and then an example mapping is given for DOCSIS and some of the other home networking architectures. We conclude with the evolution of the home networks toward a unified QoS architecture.

DOCSIS

Cable Modem (CM) and Cable Modem Termination System (CMTS) are the main entities in the DOCSIS network protocol. Several CMs- residing in the customer premises - are connected to a CMTS - residing at the cable operators' Head-End - through a Hybrid Fiber-Coax (HFC) network in a 'tree-like' structure - where the Head-End is the root and the CMs are at the leaves.

The DOCSIS downstream (DS) channel - carrying information from the CMTS to the CMs- uses a typical TV channel (6MHz wide

in USA, 8MHz wide in Europe) to carry 'Ethernet-like' packets over a continuous digital MPEG stream. As the DS channel is shared, the channel bandwidth (about 40 Mbps in USA) is distributed among all active CMs. Note that all the DS packets are received by all connected CMs. The Ethernet address is used by the CMs to filter out the packets they need.

The DOCSIS upstream (US) channel carries 'Ethernet-like' packets and uses the lower frequencies (below the range allocated to TV channels), a range that is prone to occasional interference. To cope with such situations, the US channel configuration (bandwidth, rate, error correction, and other transmission parameters) is diverse and is dynamically controlled by the CMTS. The transmission scheme is decided on a burst-by-burst basis. Again, the US channel is shared among several CMs, hence a multiple access mechanism is implemented and the channel bandwidth (up to 10Mbps in the current specification) is managed and allocated to active CMs by the CMTS. Note that US packets are received by the CMTS only. Packets from CM to CM always pass through the CMTS.

A DOCSIS domain may include several downstream and upstream channels paired accordingly to achieve the required network balance. CMs may be instructed by the CMTS to move from channel to channel as a load balancing implementation or as a means to overcome channel quality problems. A CM acts as a transparent bridge, it forwards packets that are received from the CMTS toward its local network interface (Ethernet, USB, etc') and vice-versa. Packets that are destined to the CM (such as packets to the SNMP, DHCP or other IP based agents residing in the CM) are consumed by the CM and not forwarded.

As CMs reside in varying distances from the CMTS, The Time Division Multiple

Access (TDMA) scheme implemented in the US requires subtle synchronization mechanisms. All CMs align to the CMTS clock (this clock is distributed through dedicated DS control messages). A 'Ranging' mechanism (in which the CMTS instructs the CM on the time shift and power level to use in US transmissions) is constantly active for every connected CM. This mechanism ensures that all CMs transmissions are aligned to a time base controlled by the CMTS, and that all signals are received at the CMTS at approximately the same level, ensuring the ability of the CMTS to identify collisions.

The US channel is divided into time slots. A transmission interval is a group of continuous time slots. The CMTS allocates transmission intervals for different needs (transmission requests, packet transmission, ranging messages, etc.), and transmits the allocation to the CMs using a dedicated control message in the DS. Some of the allocated intervals are multicast (generally, transmission request intervals are not allocated to a specific CM) and may result in a contention, and some of the intervals are unicast (such as a packet transmission interval). The US interval allocation messages ('MAP' messages) are transmitted in the DS by the CMTS in a timely manner. To illustrate the US transmission mechanism, consider a CM wishing to transmit a packet in the US. The CM will analyze the MAP messages in the DS until a multicast 'request' interval is allocated by the CMTS. The CM will transmit its request in the specified interval (the request contains the required transmission length) and will wait until a MAP message containing the requested allocation (a 'grant') is received. Once the grant is received, the CM transmits the packet in the allocated interval. It might be the case that more than one CM transmitted a request in the same interval. In that case, a back-off algorithm is implemented to solve the contention. To overcome the possible contention of subsequent requests, a CM may

transmit a request embedded ('piggybacked') in a packet transmission.

While the best-effort scheme that was provided by DOCSIS 1.0 was sufficient for basic Internet access, it fell short of the needs of more sophisticated services that are unable to operate in the absence of guaranteed QoS. The QoS framework of DOCSIS1.1 (as detailed below) was targeted at exactly those types of services.

Four main service categories are supported by DOCSIS1.1: Unsolicited Grant Service (UGS), Real-Time Polling Service (rtPS), Non-Real-Time Polling Service (nrtPS), and Best Effort (BE) service.

In a UGS flow, the CM is assured to receive from the CMTS fixed size grants at periodic intervals without the need to explicitly send requests. In addition to the grant size and the period, the tolerated grant jitter is also negotiated at service setup. The main advantage in using a UGS is the reduced latency achieved by eliminating the need to go through the request-grant cycle for every packet. However, using a UGS is inefficient for applications that don't require a constant data rate over time. A flavor of UGS - Unsolicited Grant Service with Activity Detection (UGS-AD) - is targeted at those exact applications (e.g. Voice with silence detection). In a UGS-AD, once the CMTS detects flow inactivity (through non usage of grants by the CM), it starts sending unicast request opportunities (also called 'Polls') at periodic time intervals. The CM can use the unicast requests opportunities to send requests (once the flow is to resume) avoiding the latency incurred by contention at multicast request intervals.

In a rtPS flow, the CM is assured to receive from the CMTS unicast request opportunities at periodic intervals. If the CM does not use the request opportunities, the CMTS allocates the reserved bandwidth to other flows, overcoming the inefficiency of UGS. In a

nrtPS flow, the bandwidth is not guaranteed to the flow. The CM, however, is allowed to use multicast request opportunities for the flow as well. The last QoS category, Best Effort, defines the minimum traffic rate (which the CMTS must reserve) and the maximum allowed rate as the main service parameters.

Multiple data flows (each flow corresponding to a service and identified by a Service ID (SID) may concurrently exist in a CM. A transmission request in the US and the corresponding grant includes the SID as the flow identifier. The CM and the CMTS negotiate the QoS for each flow upon allocation and dynamically as the service requirement changes (in dedicated procedures). The QoS is then achieved by the implementation of sophisticated scheduling mechanisms in the CMTS. A classification function is applied to every packet. The flow in which a packet is transmitted is based on the content of the Ethernet and IP header fields (allowing every application to receive a different service flow). The classification function may also indicate the suppression of the packet header (a mechanism that is useful for packets with semi-constant headers and short content, such as voice packets).

802.11

IEEE 802.11 are a set of standards for wireless networking. The original 802.11 specification, defined a network that works in the 2.4Ghz frequency band and is capable of transmission up to 2Mbps. The 802.11b standard that is becoming a dominant force in the wireless home networking arena extends this capability to 11Mbps thus providing the ability to transfer high quality audio and video content over the network. 802.11b PC cards, PCI cards and external base stations that are sold today enable multiple PCs to connect with the home or office as well as other IAs such as PDAs. The 802.11a standard provides even higher data rates (>50Mbps) using the

5.7Ghz band. The 802.11g specification will provide a direct extension to the 802.11b using the same 2.4Ghz band and providing data rates of 22Mbps.

An 802.11 network is usually composed of multiple stations and a single access point that coordinates the network activity. In some cases an Adhoc network can be established with stations only.

The original 802.11 specification used a basic contention-based MAC and as such does not provide real mechanisms for QoS (similar to plain old Ethernet). The ongoing work on the 802.11e specification is intended to provide the necessary extensions to the 802.11 MAC in order to provide real QoS. It is important to note that the original 802.11 MAC spec took in to consideration the possibility of MAC extensions. Therefore, old 802.11 devices with the original MAC will not degrade the QoS capabilities of newer 802.11e devices. However, this is not the case for Ethernet as well as Home Networking Phone Alliance (HPNA).

The 802.11e specification includes multiple levels of QoS. The basic level provides for a simple prioritization mechanism similar to the basic HPNA 2.0 mechanism and 802.3p. The higher levels of 802.11e are targeted to provide real QoS with guaranteed bandwidth and delay on a per-stream basis. These capabilities are similar and, in some cases, exceed the ones that exist in the DOCSIS 1.1 specification.

The 802.11e specification includes multiple bandwidth management schemes. The lower levels employ a distributed approach where each station on the network is making its own QoS decisions similar to Ethernet. The highest level (3) uses a centralized approach where the access point controls the assignment of bandwidth to the client units similar to a DOCSIS CMTS and DOCSIS CMs. All these schemes can coexist on the same physical channel. The coexistence

of different level QoS devices is similar to the co-existence capability of DOCSIS 1.0 and DOCSIS 1.1 modems on the DOCSIS network.

The 802.11e draft defines a set of MAC sub-layer QoS parameters more extensive than may be needed, or may be available, for any particular instance of QoS traffic. These parameters, collectively called a traffic specification and applied to a traffic category (TC), are Traffic Type, Ack Policy, Delivery Priority, Retry Interval, Polling Interval, Transmit Interval, Nominal MAC Service Data Unit (MSDU) Size, Minimum Data Rate, Mean Data Rate, Maximum Burst Size, Delay Bound, and Jitter Bound.

HPNA

The network defined by the HPNA enables devices to transmit packets over residential phone lines. Two specifications have been developed: HPNA 1.0, which offers a data rate of 1 Mbps using an Ethernet-like Carrier Sense Multiple Access (CSMA) protocol without any QoS mechanisms, and HPNA 2.0, which offers a data rate of 10 Mbps using a prioritized CSMA protocol. In the discussion that follows, we refer to the later HPNA 2.0 protocol.

The management of the PHY and MAC layers is distributed. All the devices cooperate to enable minimal overhead of network resources usage for network management. All decisions, such as transmission schemes and priority, are implemented locally in every device. For example, every device decides, on which transmission scheme to use based on the destination device(s) characteristics.

The network supports eight priority levels from zero (lowest) to seven (highest). A time-division scheme is used to enable prioritized access. For every priority level, a designated time slot is defined (the time slot for priority seven being first, followed by the time slots

for the other priorities in descending order). Devices are allowed to start the transmission of a packet only in (or after) the time slot that corresponds to the packets' priority (as decided by the device). Since collisions can still occur - usually between packets of the same priority - a contention resolution algorithm (CRA) - based on a random back-off mechanism - is implemented. New packets for transmission do not preempt the CRA unless they are of higher priority.

The above scheme carries some inherent characteristics that impact the usability of the HPNA network for services requiring guaranteed QoS. Ignoring collisions, the worst-case latency of a packet to be transmitted is in the range of the maximum transmission time of a packet - about 4.2 milliseconds. If collisions occur, however, the latency can be only statistically bounded (every contention resolution 'round' takes about 200microseconds). As the stations are free to choose the transmission priority of most of the packets (although a mapping of IEEE 802.1D priority levels to HPNA priorities is provided), the resulting network latency can't be guaranteed (higher priority 'preempting' packets may affect the CRA completion time).

The above limitations call for higher-level network synchronization - a management entity at a higher layer that controls both the priority and timing of packet transmission at the MAC and PHY layers. Such a scheme will enable the provision of quality demanding services such as voice and video over phone lines.

IEEE1394

The IEEE1394a standard defines a wired networking targeted to connect very-high-speed multimedia devices over short distances. The IEEE1394a supports connection speeds of up to 1.6Gbps with lower levels starting at 200Mbps. IEEE1394

uses special types of cables to carry the information. Multiple IEEE1394 devices may be connected to each other in a daisy chain formation. IEEE1394 interfaces are part of every digital camcorders and are also widely deployed in multimedia PCs. It is also used to connect high-end video display devices (HDTV) and editing equipment. The IEEE1394 was also chosen as the interface between the next generation of OpenCable™ set top boxes and HDTV displays. Efforts are being made to extend the capabilities of IEEE1394 and to provide for wireless variations. Other IEEE1394 related specifications define the way MPEG-2 TS (the standard used for Digital broadcast over cable and satellite) are carried over IEEE1394 networks, delivery of IP traffic over IEEE1394 and control language over video equipment (start, play, stop, forward, back, titles etc.).

The IEEE1394 MAC layer includes two types of channels: asynchronous and isochronous.

The asynchronous channel is a best-effort data channel that has no QoS mechanism associated with it.

The isochronous channel is a reservation-based, assured-bandwidth channel.

A designated isochronous resource manager in one of the devices connected to the IEEE1394 network manages the Isochronous channel. The resource manager accepts reservation requests from all devices on the network and assigns specific bandwidth to each one of them. The reservation should reflect the peak rate of transmission from the source device. Bandwidth that is not assigned to any Isochronous reservations or is not utilized by an existing reservation is used to carry best-effort asynchronous transmissions.

It is important to note that the capacity of the DOCSIS channel is negligible compared to the capacity of 1394, which is typically 200-400Mbps and the garnered capacity to a subscriber using DOCSIS channel is a fraction

of the overall 40Mbps D/S channel). Also the jitter and delay provided by IEEE1394 on the isochronous channel is extremely low due to the high bit rates and fixed reservation mechanism.

BlueTooth

The Bluetooth (BT) specifications define a short-range wireless network working on the 2.4Ghz band. The maximum data rate supported on a BT system is 768Kbs. The BT specification defines two connection methods: The point-to-point connection where two devices are creating a private connection, and the point-to-multi-point connection where a real network of up to eight active devices is created. The point to multi-point connection is based on an assignment of a master device to each such network (named a piconet) and multiple slave devices. A slave device may belong to multiple piconets. BT is rapidly deploying today in devices such as cell phone PDA etc..

The definition of mapping in this section relates to the point-to-multi-point connection defined in BT.

BT defines two types of channels for transmission - synchronous and asynchronous. Synchronous channels (SCO) are symmetric and provide a 64kb/s bi-directional connection between the Master and a specific slave. Transmission and receiving slots are sent periodically with a fixed interval between them, each slot is 625usec in length. Up to three such slots can be accommodated by a single master or slave. The SCO slots are setup by the master using the BT LM protocol. There are multiple types of SCO slots, most of them are targeted to voice distribution, but one of them can carry both voice and data simultaneously.

Asynchronous packets (ACN) are sent on the slots left after SCO assignment. The slaves send information only after they receive

information from the master. There are multiple types of ACN slots that differ by their payload size and FEC protection. For the ACN channels, the BT specification defines the L2CAP layer that enables segmentation and re-assembly of packets as well as QoS services and connection establishment.

The SCO channels of BT are really targeted at voice streams for which they provide QoS. The different types of slot VH1..3 provide different packetization periods and delays for the voice transmission. It is important to note that these slots are highly tailored for simple PCM-based transmission of voice packets. Usage of these slots for other applications or voice codec types is not trivial. The usage of SCO channels is also limited since the BT L2CAP conversion layer is not defined for this channel and thus the direct translation to IP is not obvious. It is most likely that SCO channels will be used for voice delivery only.

Best-effort data as well as reservations not using SCO will use the ACN slots. The L2CAP layer of BT runs on top of ACN channels and includes a QoS reservation protocol that uses values, such as bucket rate and jitter, to define the QoS of a designated stream. The L2CAP QoS is not mandatory and so not all BT implementations support it.

From the QoS perspective, BT SCO channels provide for excellent QoS but are intended for a very narrow set of voice applications. The L2CAP QoS mechanism can potentially provide for good QoS but it is not mandatory as part of the spec. The overall bandwidth capabilities of BT limit its ability to support high-end applications such as mid/high quality video streaming, as such it can extend DOCSIS QoS for a limited set of applications.

HomeRF

The HomeRF specification is another standard for wireless networking. HomeRF devices are capable of transmitting up to rates of 1.6Mbps. Future HomeRF specifications are targeted to reach the 10Mbps range. HomeRF devices marked the first phase of commercial wireless home-network devices. HomeRF PC Cards, PCI cards and external modems that are being sold today enable the connection of multiple PCs and peripherals with in the home.

The HomeRF MAC layer architecture is a combination of an asynchronous, Ethernet-like access mechanism and an isochronous, circuit-switched TDMA access mechanism.

The isochronous channels are used primarily for up to eight active 32 kbps ADPC voice connections using DECT signaling in the upper layers. As in BT SCO channels, these channels are providing full QoS with guaranteed delivery and bounded jitter and delay, but only for very narrow range of voice applications.

The asynchronous data service is comprised of both a best-effort service and a prioritized asynchronous service. The management of the PHY and MAC layers is distributed. All the devices cooperate to enable minimal overhead of network resources usage for network management. All decisions, such as transmission schemes and priority, are implemented locally in each device

HomeRF devices support eight prioritization levels for asynchronous data as well as a lower level of best-effort data.

Specific streams can be assigned higher priority over other streams. It is also possible to define a nominal time reservation and maximal time reservation for a stream.

Like 802.1p, the HomeRF asynchronous channel can only provide for prioritization but with no guarantee of delay or bandwidth.

Its limitations are similar to the ones identified for HPNA 2.0.

The current spec of HomeRF is also limited in the application it can carry due to its bandwidth. It would not be suitable to carry mid/high quality video streams.

HomePlug

Although power-line, as a communications medium, is more complex than wireless and phone-line, the distribution of outlets within the home makes it a viable platform for home networking. The HomePlug Powerline Alliance was formed to create an industry standard for high-speed home networking over power lines. Similar to the situation with phone-lines, the management of the PHY and MAC layers is distributed - all decisions are implemented locally in each device depending on a dynamic configuration that is suited to the specific pair of interconnected outlets. The adaptation to specific channel characteristics is essential to the efficient use of the media since the exact channel behavior is unique and generally unpredictable. An encryption scheme in the MAC layer copes with the inability of the distribution transformer to block the propagation of power-line signals from one home it powers to another.

A variant of the Carrier Sense Multiple Access with Collision Avoidance (CSMA/CA) shared access protocol is used (with a random back-off algorithm for contention resolution). QoS is supported through the additional definitions and mechanisms for priority classes and latency control: stations use a priority resolution interval to signal the requested transmission priority, ensuring the early transmission of the higher priority frames. Latency is reduced through the use of segmentation and reassembly (a mechanism is provided to allow uninterrupted transmission of multiple segments in the absence of higher priority

frames). By enforcing the discard of stale packets, a bound on the latency is achieved. This is particularly important for applications demanding latency control (e.g. VOIP).

HomePlug shares many of the phone-line protocol limitations, primarily the inability to prevent a non-cooperating - though compliant-device with 'private' priority semantics. To truly support latency and jitter sensitive applications, higher layer coordination is a desired scheme in this case as well

QoS Based on flow-specs requirements

We are interested in the bridging of QoS from the DOCSIS technology to other home-networking technologies. This raises a question regarding what is actually being bridged. As can be seen from the previous sections of this paper, there are multiple ways to achieve QoS over different networks technologies, some use CBR mechanisms (Constant Bit Rate) where a set of two numbers (Packet rate, Packet Size) define the QoS, while others use prioritization mechanism where a single, different value (Priority) defines the QoS. Other mechanisms use even more complex representations for the QoS with multiple values. Trying to directly match these QoS mechanism and values is in most cases impossible.

A different approach to QoS mapping would be to look at the actual streams of data for which we are attempting to provide QoS. We can use some sort of general QoS requirement definition for these streams, and then try and map these general requirements to any one of the different technologies on the stream's way from source to destination. The mapping of two technologies is indirectly performed by mapping each to these general stream requirements.

One general model that had been used in multiple standards is the flow model, which imitates a fluid-flow model. In this model, the

data flows are represented as buckets of fluid (the data is the fluid) that are poured into pipes of variable size. The flow is defined by the size of the bucket, the number of buckets per second, the peak rate for which a fluid can pour from the bucket and the pipe size. Additional values such as the over all delay can be added to provide for more information on the requirements of the stream. Although it seems very crude, the flow model has been proven to be a viable tool to defining the real QoS requirements of data streams. These general flow requirements are titled “flow spec.”

The benefits of this indirect mapping are not obvious at first glance. In order to better explain them let’s consider the following example:

Technology A uses a QoS mechanism that is based on the assignment of a constant bit rate (CBR) channel for each stream. The channel is defined by the packet size transmitted and the time-gap between each of these packets, which is completely fixed.

Technology B uses a prioritization mechanism with limitation on packet sizes. Streams with higher priority are given access to the channel before streams with lower priority. The maximum length of a packet to be transmitted for each stream is also limited.

We wish to deliver the following stream that is defined by the flow model as follows:

Bucket Rate = 100 per second

Bucket Size = 100 bytes

Peak Rate = 10Kbytes/per second

Pipe bandwidth = 10Kbytes/per second

Max Delay = 30msec

Technology A may map these stream requirements to a CBR channel of packet size

200 bytes and repetition of 50 times per second. This definition satisfies the requirements of the stream.

Technology B would have to consider the other existing streams on the network and consider their relative priority and packet sizes. Let’s assume that the network data rate is 100K byte per second, the number of priority levels is four, and three streams of level three already exist with a packet size of 2000 bytes each. The new stream will be mapped to level four with a packet size of 100 bytes. Mapping it to lower levels will cause a violation of the Max Delay defined by the stream.

We were able to map the stream to both technologies and, by that, create an indirect mapping between two completely different technologies. This mapping will enable us to deliver the specific stream over a network that is comprised of segments from both of these technologies while maintaining the stream requirements. However, we cannot assume that the direct mapping exists between the two. A “technology A” CBR channel of 200bytes and repetition of 50hz is not equivalent to “technology B” level four. The mapping exists only in the context of the specific streams that are active in the system on a specific time.

Mapping to Flow Specs

This section provides a short explanation of the mapping between flowspec parameters and some of the actual technologies previously described.

DOCSIS

The DOCSIS QoS is using multiple types of services such as CBR, Real Time Pole (RTP) etc.. The information provided by the flowspec can be used by the DOCSIS CMTS

to decide what type of service a specific stream will receive. Streams that have their peak rate equal or close to their average rate will receive a CBR service and the parameters of this CBR will be set according to the peak rate of the flow spec, other streams that have a more bursty nature may use an RTP service. A more detailed mapping between DOCSIS parameters and flow specs can be found in CableLabs™ PacketCable™ QoS specification. It is important to note that the mapping information provided for DOCSIS is just a guideline for implementation. Every CMTS vendor can decide on different mapping algorithms. This does not create any interoperability or performance problems because DOCSIS uses a fully centralized bandwidth and QoS management.

802.11

The QoS mechanisms defined in 802.11e use parameters that are very similar to the ones of the flow specs. Therefore, a direct mapping between the two can be made. The physical mapping of these requirements to the channel behavior depends on which level of 802.11e QoS is used. Some of the levels use a distributed approach where each station is making its own decisions on the scheduling of packet for transmissions, while others use a centralized approach where a master station makes the decision for all other stations. In general the centralized approach can provide full compliance to the flow-spec requirements, while the distributed approach is usually statistical, and can provide a very high chance of compliance to the requirements, but not complete certainty.

In most cases, an 802.11b/e network can meet flow-specs requirements that can be met by the DOCSIS channel.

BT

There are multiple options to providing QoS over BT. In some cases, the flow-spec

requirements can be satisfied by using an SCO channel (given the stream bucket size and rate fit the size of the SCO slots), other flow specs may be mapped to the L2CAP QoS parameters. These parameters are very similar to the flow-spec parameters and a direct mapping is possible. The scheduler in the master station performs the physical mapping of these requirements to the channel behavior. It is important to note that, due to the limited bandwidth of the BT network, it is likely that some streams flow specs that the DOCSIS network can satisfy would not be met by the BT network (i.e. a flow spec of data rate higher than 768kbit/sec). For BT devices that do not include the L2CAP QoS extensions, it is generally impossible to assure the fulfillment of the flow-spec requirements. If the rates of data in the flow spec as well as the delay requirements are significantly lower than the overall BT channel capacity, then there is a good chance that the flow-spec requirements will be met. But for rates that are closer to the channel capacity or strict delay requirements, the chances are that the non QoS BT network would not be able to satisfy these requirements.

HPNA

The current HPNA specification (2.0) uses a priority based QoS mechanism. The exact usage of the priority mechanism is not defined in the HPNA specification and is left open. Flow specs with higher bandwidth and delay requirements may be mapped to higher priority while best-effort or low-capacity/high-delay flow specs will be mapped to lower priorities. This type of mapping can guarantee the requirements of the flow specs when we are dealing with a low number of streams. When the number of streams increases, we have a situation where multiple streams use the same priority level and it is impossible to guarantee the delay that packets from each of these streams may encounter. Using statistical models we can in most cases

assure to a very high probability that the flowspec requirements will be met even in cases where multiple streams are using the same priority.

Adding a control protocol that will negotiate scheduling information between the different devices on the same HPNA network can make an improvement to the QoS provided by the basic prioritization mechanism. Such protocols are currently under investigation in different standardization groups.

It is important to note that all the described mechanism for QoS over HPNA require that all devices on the network will obey a defined set of rules for the usage of the priority mechanism. These rules are not part of the HPNA specification itself and since the scheduling of packets transmission is distributed (each HPNA device makes its own decision on when to transmit, there is no central point of coordination), it is hard to assure the behavior of a network that is using equipment from multiple vendors. The most obvious example of this problem is the connection of an HPNA adaptor to a PC where the priority level is set by the PC owner to the maximum (7), and large packets are sent at a high rate over this interface. Even if all other devices on this HPNA network are attempting to cleverly schedule streams to priorities, they will all be blocked by the transmission originating from the PC.

HomeRF

The HomeRF QoS mechanism is a combination of CBR channels and a priority based QoS channel. Similarly to BT, flow specs that fit into the HomeRF CBR channels may use it for delivery of streams (very narrow usage for voice streams). Other flowspecs will need to be mapped to the different priorities. The problems associated with mapping to a priority based QoS are

explained in the HPNA section. These problems are significantly more severe when dealing with HomeRF as the capacity of the HomeRF channel is much lower than the one offered by HPNA 2.0. As in the case of BT, some of the flowspecs delivered by DOCSIS would not be deliverable over HomeRF because of its limited bandwidth.

IEEE1394

The IEEE1394 QoS is based on assignment of CBR channels for each stream. The mapping from flowspec requirements to CBR can be based on the R (pipe bandwidth) parameter of the flowspec. This is a very crude and simple mapping that does not take into account the other flow-spec parameters. But, for all practical applications, the other requirements of the flow spec will be met since the IEEE1394 channel overall bandwidth is large and the delays are short. This is specifically true when talking about streams that need to be carried over the DOCSIS channel that offer smaller bandwidth and longer delays. In general, streams that are carried over DOCSIS networks should not have problems being carried over IEEE1394 networks.

Conclusion:

The evolution toward a unified framework

The current home networks are mainly used to share Internet connection and a printer. However, they serve as the baseline for the future home networks. As more applications, appliances and services become available, the underlying components of the

future home network will be aggregated, resulting in a heterogeneous home network.

In today's home network, the home appliances, such as TVs, phones, PCs and, PS2, are generally owned and controlled by the home resident. Some of those end points are 'general purpose' in nature, meaning they are not dedicated to a specific service. Rather they may be used for different services upon need. Some others are dedicated to a service and, as such, are more prone to end up being owned and controlled by the service provider. The home network itself, connecting the end devices to the service sources and making the services possible, is a clear candidate to be handed-over (in terms of management and control, rather than ownership) to the responsibility of an entity other than the home resident. The diversity of configurations will make it impossible for the home resident to support and maintain it by itself.

A unified framework is thus required to enable external control of the home network and a way to analyze the network topology, its components and their capabilities. For every network component, a method is required to control and configure it to other components in order to enable the service provision and assure the quality of service end-to-end. The components of such a framework are being defined today in several standardization efforts (OSGI, UPNP, CableHome et al.). Some of those standards complement each other, and some compete. Only a set of standards that solves all the aspects of controlling the home network (QoS, security, management and configuration, etc.) will prevail.

The CableHome™ specifications defined by CableLabs, aim at precisely that. The concept of generic flow specs lies at the heart of the CableHome model, so as to incorporate IP- based mechanisms such as SNMP, SBM, RSVP and others. However, to prevail as the 'end of the road' framework (and this applies to other candidate frameworks as well) special

attention must be given to supporting an evolutionary process. Both existing components and those to be aggregated to home networks while the frameworks converge will have to be supported.

Allowing an evolution rather than imposing a revolution is the main challenge of the future home networking framework.

CABLE TELEPHONY PERFORMANCE OVER DWDM NETWORKS

Ricardo A. Villa, Senior Engineer, DWDM Systems Engineering
Sudhesh Mysore, PhD., Director, DWDM Systems Engineering
Oleh Sniezko, Vice President, Engineering

AT&T Broadband, 188 Inverness Drive West, 3rd Floor,
Englewood, Colorado, 80112

Abstract

In the last two decades, a trend toward metro market clustering under a single operator led to changes in metro market architecture. This trend allowed for significant headend consolidation and for lowering the number of signal processing facilities to the level dictated by local programming requirements and operational issues. A side effect of this benefit was an increase in distances between the processing facilities and the customers. This had not been an issue for traditional broadcast services that did not require two-way communication. However, with the introduction of interactive services with their strict requirements for proper timing and synchronization, the absolute distance from the processing centers and the differential distance between the devices transmitting on the upstream path became critical.

Additionally, requirements for increased network availability on long fiber routes made it necessary to deploy optical switching to reduce the network downtime to a minimum. Optical switching activation results in an instantaneous change in distance that causes a change in delay between the terminal devices affected by the switch and the processing facilities. It also affects differential distances between the terminal devices. Therefore, a reliable and timely recovery from these changes became crucial to achieve network availability targets.

This paper presents the requirements formulated for all interactive service platforms operating in the HFC network with consolidated signal-processing facilities. It also presents and discusses the results of tests performed on several proprietary digital telephony platforms to verify these requirements.

INTRODUCTION

The major system clustering and headend consolidation activity occurred between 1990 and 1997. It brought a significant benefit in operational savings, improved signal quality and network reliability. In 1994 and 95, a deployment of highly interactive services began. Initially, for high-speed Internet access services, the segmentation requirements were low enough to allow for a cost-effective use of 1310 nm technology for short distances and 1550 nm externally modulated laser transmitters for longer distances. However, with the deployment of digital telephony services and increased segmentation requirements, expensive 1550 nm externally modulated laser transmitters dedicated to a small number of customers could not be justified.

Additionally, increased requirements for network availability made it necessary to deploy self-healing rings deeper into the network. This goal could be achieved only after a reduction in fiber count in these rings.

Primary Ring
60,000 - 200,000 Homes Passed

10,000 - 40,000 Homes Passed

Legend:
 □ Headend
 PH Primary Hub
 □ Node

Initially, due to the lack of other cost-effective transport technology choices, signal processing equipment for high-speed Internet access services (cable modem termination systems — CMTSSs) and for digital telephony (host digital terminals — HDTs) was deployed in secondary hubs and interconnected with the primary hub rings with SONET and ATM transport systems¹. In many cases, point-to-point Ethernet transport systems were deployed in place of ATM systems to support high-speed Internet access equipment. This solution is depicted in Figure 2. Although this negated the trend toward, and hence the benefits of, processing facility consolidation, it was the only viable solution at that time. In this arrangement, redundancy switching or routing was taking place on the network side of the processing equipment with well-defined timing performance of the SONET, ATM or Ethernet transport systems. The CMTSSs and HDTs were connected over an HFC network to their corresponding cable modems and network interface units (NIUs)

The diagram illustrates the system architecture, divided into three main functional areas: HEADEND, HUB, and a Broadband section.

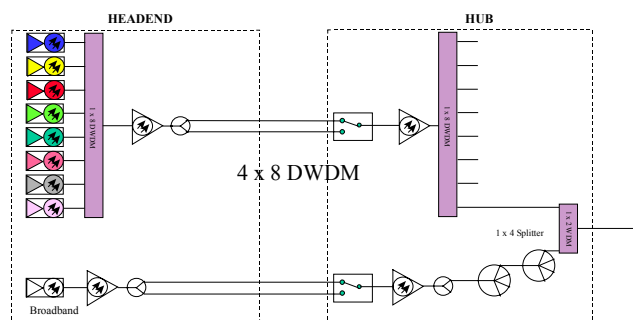
- HEADEND:** This section includes an "Add Insert" block, a "Remux-Scheduler for SD-WAN" block, and a "SONET Bay" block. It is connected to the HUB via an 8 dB link at 1550 nm and a 13 dB link at 1310 nm.
- HUB:** This central section contains a "SONET Bay" block, a "Voice Bandwidth Manager and Modem" block, a "Data Modem & Router" block, and a "Combiner / Splitter" block. It is connected to the HEADEND and the Broadband section via 5 dB links at 1550 nm and 7 dB links at 1310 nm.
- Broadband Section:** This section includes a "Broadband" input, a "Combiner / Splitter" block, and a "Node Single Rx" block. It is connected to the HUB via a 7 dB link at 1310 nm.

The diagram shows the flow of data and voice signals between these components, including links at 1550 nm and 1310 nm.

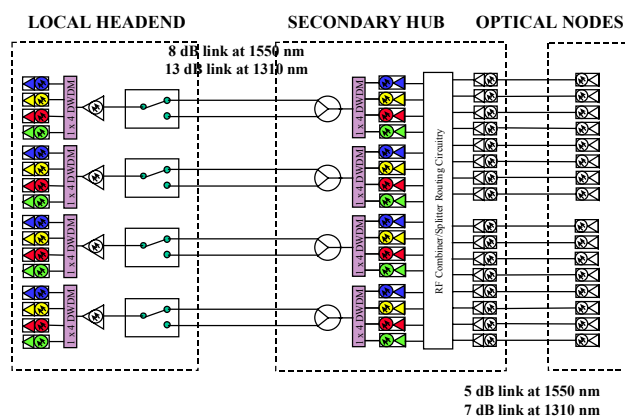
The diagram illustrates the system architecture, divided into a HEADEND and a HUB. The HEADEND contains a 100 BaseT transceiver. The HUB contains a 100 BaseT transceiver, a Data Modem & Router, and two Combiner Splitters. A red double-headed arrow indicates a link between the two 100 BaseT transceivers, labeled '8 dB link at 1550 nm' and '13 dB link at 1310 nm'. The HUB also has a 5 dB link at 1550 nm and a 7 dB link at 1310 nm. A Broadband input is shown at the bottom left, connected to the HEADEND. A Node Single Rx is shown at the bottom right, connected to the HUB via a 2x2 coupler.

Figure 2: Hybrid Analog and Digital

In 1998, dense wavelength division multiplexing (DWDM) technology reached a level of performance adequate for carrying multiple QAM signals over long distances. Moreover, it reached the price points to become a cost effective alternative to the SONET transport system in secondary hub rings². It also eliminated the need for different transport systems dedicated to different services. A DWDM network configuration is presented in Figure 3.



a) Forward



b) Reverse

Figure 3: DWDM in Secondary Hub Links

The most important benefit of this technology was the fact that the operators could continue the consolidation of signal processing facilities. The CMTSS and HDTs could again be concentrated in a few primary hubs (PHs) and headends. This resulted in cost savings in hub buildings, improved reliability due to faster mean time to repair (MTTR), and operational efficiencies due to having highly sophisticated and environmentally sensitive equipment placed in fewer, well-designed locations.

This desirable change increased the distances to the cable modems and NIUs from the typical 12-30 km to much greater distances of 70-100 km and, in a few extreme cases in backup routes in secondary hub rings, to 150 km. Due to these

increased distances, it also became necessary to introduce diverse fiber path routing between CMTSSs and HDTs and their respective cable modems and NIUs, both in the forward (FWD) and reverse (REV) paths. The DWDM optical links between PH and SHs were configured as self-healing rings. Each SH was connected with its respective PH via automatic optical protection switches (AOPSs) selecting between principal and backup fiber routes, often significantly different in length (and hence in transit delay). These differences could result in instantaneous changes in delay and/or optical power levels as a result of the switchover. The latter can be addressed by a judicious use of optical amplifier and attenuators for level equalization in principal and backup paths after redundancy switching. However, the instantaneous changes in delay caused by AOPS activation require cable modems and NIUs to adapt to such changes with minimal loss in service availability.

The DOCSIS standard addressed the long absolute and differential distances likely to exist between CMTSSs and cable modems, and DOCSIS-compliant high-speed Internet access platforms can operate at these long distances (transit delay from headend to most distant customer can be ≤ 0.800 ms). However, the proprietary digital telephony platforms could not initially meet similar requirements for absolute and differential distances. Moreover, neither DOCSIS standards nor proprietary digital telephony platform specifications directly addressed the recovery issues related to the instantaneous changes in these distances.

This situation necessitated the definition of new requirements for the interactive service platform capabilities for both distance requirements and recovery

requirements in case of the instantaneous change in these distances.

REQUIREMENTS FOR INTERACTIVE SERVICE PLATFORMS

AT&T Broadband Requirements

Based on the considerations presented above, AT&T Broadband defined the following architectural requirements for the interactive service platforms:

1. Absolute maximal distance between signal processing equipment and terminal equipment at which the terminal equipment can range and register shall be as specified below. All terminal devices must register automatically for the entire absolute range without reconfiguration of the headend/primary hub equipment.

- **Requirement: 70 miles,**
- Objective: 90 miles.

(**Note:** DOCSIS requirement approximates to 110 miles of optical fiber)

2. Differential distance between the closest and the furthest terminal devices served by a single headend modem (i.e., on the same forward and upstream carrier path):

- **Requirement: Half of the maximal absolute distance,**
- Objective: Maximal absolute distance.

3. Change in the distance due to the optical path redundancy switch activation — all terminal devices at the maximal unit load per carrier path, calculated based on traffic modeling and equipment limitations, must re-range and re-register after a change in the distance for either or both carrier paths (forward or reverse) equal to the maximal absolute distance

in either direction (i.e., from the maximal distance to zero distance and from zero distance to the maximal distance):

- **Requirement: Within 4 minutes,**
- Objective: Within 2 minutes.

(**Note:** The additional objective is not to terminate a call in progress upon the change in distance described above).

4. Terminal units operating on the same modem as the units subjected to the change in distance (optical path redundancy switch activation) shall not require re-ranging or re-registration or be affected in any way.

Statistical Analysis of PH to Home Distances – Requirement 1

Approximately 500 SHs are designed or scheduled for design and construction in AT&T Broadband markets with interactive services. Approximately 5% of these hubs (25 locations) are located far from a primary hub and either a principal or backup route or both would result in exceeding a distance of 50 miles between the primary hub and the terminal equipment (cable modem or NIU). Some of the older versions of the digital telephony platforms would not operate at these distances and would require placement of HDTs in secondary hubs. This would result in a significant capital cost and an increased cost of operations. This consideration led to the requirement for longer operational distances. Several secondary hubs are located further than 70 miles away from their respective primary hubs. The distance from any primary hub to the furthest customer in the existing AT&T Broadband markets does not exceed 90 miles.

Clustering – Requirement 2

In extreme cases where some terminal devices fed via home run nodes (nodes fed directly from primary hubs) are attached to the same modem as the terminal devices fed via secondary hubs, the differential distance between the devices and the signal processing equipment may approach the absolute distance to the furthest customer. Although the network designers and service operators try to avoid these situations, an efficient use of the signal processing equipment dictates the need for this requirement.

Network Availability – System Recovery: Requirements 3 and 4

The optical fiber links stemming from primary hubs can reach lengths approaching 90 miles. Based on the data collected over the years³, this length of unprotected fiber link could result in 107 minutes of downtime a year. Hence, the redundancy switching in these and any other optical fiber runs, except for the direct runs from primary and secondary hubs to the nodes (usually shorter than 16 miles, and typically only several miles), is required to meet the network availability goals. Since the SONET or ATM rings to the SHs would prevent any downtime caused by fiber cuts, the downtime allocations for fiber-cut redundancy switching in these links was equal to zero. Therefore, AT&T Broadband attempted to reduce to a minimum the redundancy switching downtime caused by re-ranging and re-registration. It was assumed that there would be one redundancy switching event (i.e., two switchovers per event) a year due to maintenance activity in addition to the redundancy switching caused by fiber cuts. Based on the statistical data on fiber cut events, there is a 40% likelihood

of a cut per year in a 90-mile optical cable run. The 4-minute requirement was selected as a reasonable and achievable compromise based on discussions with several vendors of proprietary and DOCSIS-compliant systems. However, it does not fully satisfy AT&T Broadband's needs. It would still result in an 11.2-minute downtime in longest runs versus the zero downtime for SONET rings. The objective of a 2-minute downtime with the additional objective of not dropping calls in progress would be more satisfactory.

TEST RESULTS

Requirement 1

To preserve the clarity of the test result analysis, the test setup and test conditions are detailed in the addendum at the end of this paper. The table below illustrates the results for each cable telephony platform.

Table 1: Test results for the Absolute Distance

Absolute Maximal distance at which NIUs can Register

	70-mile (112 km) Requirement	90-mile (144 km) Objective
Platform A	Yes	Yes
Platform B	Yes	Yes
Platform C	Yes	Yes

The results indicate that all three platforms performed the normal provisioning, registration, diagnostic and operations functions at the maximal required distance of 70 miles (112 km), and also at the objective distance of 90 miles (144 km) at high call volumes without problems, thus satisfying Requirement 1.

Requirement 2

The test setup and test conditions are detailed in the addendum at the end of this paper. The test results are shown below in Table 2.

Table 2: Test Results for Differential Distance

	Requirement: 1/2 Max absolute distance (35 miles, 56 km)	Objective: Max absolute distance (70 miles, 112 km)
Platform A	Yes	Yes
Platform B	Yes	Yes
Platform C	No	No

Platform C did not meet this requirement at the time of the test. This platform was also tested to make sure that it did not work *outside* of its specified differential range. An attempt to register units at distances outside of the specified differential distance range after registering units at 70 miles or other distances could result in interference and service disruption for the units already registered. These tests were performed during normal operations under a call load and showed that there was no interference and the platform did not attempt to register the units placed outside of the specified differential range.

The differential distance tests were also performed at 90 miles for all platforms with the results being identical to those in Table 2.

Requirements 3 and 4

When an optical fiber protection switch activates in the FWD or REV paths, NIUs experience an instantaneous change in delay offset from the HDT since the fiber path being switched-to is usually of a different length than the original path. Frame alignment in the FWD and REV

becomes offset in time and the NIUs no longer transmit or receive at the expected time intervals. The system detects this condition and acts to correct it. As a result, most NIUs will un-register from the HDT modem (i.e., lose communication). Usually, most of the calls in progress are lost as well. The system will then attempt to re-register the NIUs returning them to full operation within a certain amount of time. Depending on the amount of delay change, some platforms can preserve the active calls, and some or all units may not un-register but, in general, there is always a fiber delay change value for which all units un-register and all calls in progress are lost. If NIUs un-register, they are unavailable for calls until fully re-registered. This situation constitutes a real service outage. Detection “audits” exist within each system platform to detect that a loss of communication has occurred, and recovery mechanisms are triggered as a result of the audits to re-register the NIUs. Each vendor has a different implementation of these mechanisms and audits. Therefore, testing of the AT&T Broadband requirement necessitates a thorough understanding of the particular implementation to create the appropriate test conditions to verify its effectiveness.

The time to re-register depends on different factors for different platforms. The most important factors are: the number of NIUs registered on the HDT modem, the amount of delay change, the amount of traffic being handled by the HDT modem under test during the switching event, and whether delay was increased or decreased as a result of switch activation. Also, audit and recovery mechanisms are triggered by different situations in each platform. For example, one platform is more sensitive to small changes in delay and another is more sensitive if the delay is decreased rather than increased.

The test setup and test conditions are detailed in the addendum at the end of this

paper. Table 3 summarizes the test results.

Table 3: The Platform Behavior after Optical Redundancy Switch Activation (Separately for FWD and REV)

FWD Switching		0 km (no change in delay)	1km to 4 km	20 km, 25 km, 50 km, 75 km, 100 km, 112 km (70 miles), 125 km, 144 km (90 miles).	Notes
Platform A	Add Delay	No change, no un-registration	2 min 57 sec. to 4 min 53 sec.	2 min 49 sec. to 3 min 33 sec.	Fully loaded HDT modem. Most units unregistered during a switch event except for those at 0 km. Times are to full recovery.
	Reduce Delay	No change, no un-registration	3 min 17 sec. to 6 min 15sec.	2 min 45 sec. to 3 min 53 sec.	
Platform B	Add Delay	No change, no un-registration	Not tested	10 sec. to 3 min 33 sec.	Only 8 NIUs per HDT modem. Majority of units un-registered and re-registered, including ones at 0 km.
	Reduce Delay	No change, no un-registration	Not tested	3 min 17 sec. to 4 min.	
Platform C	Add Delay	No change, no un-registration	No un-registration	No un-registration (20, 25 km) to 5 min 22 sec.	Fully loaded HDT modem. If un-registration occurred, all units un-registered, but never those at 0 km.
	Reduce Delay	No change, no un-registration	No un-registration to 6 min 23 sec.	3 min 48 sec. To 5 min 47 sec.	
REV Switching		0 km	1km to 4 km	20 km, 25 km, 50 km, 75 km, 144 km (90 miles).	
Platform A	Add Delay	No change, no un-registration	Not tested	3 min 12 sec. to 3 min 17 sec.	Fully loaded HDT modem.
	Reduce Delay	No change, no un-registration	Not tested	2 min 45 sec. to 3 min 56 sec.	
Platform B	Add Delay	No change, no un-registration	Not tested	10 sec to 3 min 30 sec.	Only 8 NIUs per HDT modem, far below the full load. Also units at 0 km un-register at times.
	Reduce Delay	No change, no un-registration	Not tested	2 min 18 sec. To 3 min 43 sec.	
Platform C	Add Delay	No change, no un-registration	No un-registration	No un-registration (20, 25 km) to 5 min 5 sec.	Fully loaded HDT modem. Higher sensitivity to delay reduction than to delay increase.
	Reduce Delay	No change, no un-registration	No un-registration to 5 min 43 sec.	3 min 48 sec. to 5 min 47 sec.	

The results indicate that Platform A met the 4-minute requirement for fiber deltas of 20 km and greater. In contrast, it did not meet the requirement for fiber deltas of 1 to 4 km. This platform was the only one that had a specific system mechanism to detect and correct changes in delay caused by the fiber switchover. The other two platforms had not developed specific methods to deal with this at the time of this writing. They relied on mechanisms that were optimal for recovery from outage situations and not specifically from changes

in transit delays. The improvements to fiber switching for Platform A are implemented in software and do not require new hardware. Platforms B and C did require new hardware and new types of units to meet some of the new requirements.

Platform B showed that re-registration times were just at or below four minutes. However, this platform was loaded with only eight units per modem, far below its maximal capacity. It is reasonable to extrapolate that this time will most likely

exceed the 4-minute requirement for fully loaded modems. Other tests performed by the vendor have confirmed this extrapolation. During some of the fiber switchovers, the units at 0 km un-registered and re-registered even though their delays were not affected by the switching event. This behavior did not occur in the other two platforms. The vendor is working to implement a mechanism that specifically detects fiber delay changes and responds to them in less than four minutes, without affecting the units not subjected to the delay change (i.e., not subjected to the fiber switchover) on the same modem.

Platform C was very robust when fiber switchovers resulted in the addition of a delay caused by up to 25 km of additional fiber. In these situations, units did not un-register and calls were not dropped. However, reductions of delay by more than

1 km equivalent fiber caused complete un-registration and subsequent re-registration. In both situations (negative delay changes and positive delay change in excess of 25 km of fiber), the recovery times were often above the 4-minute requirement. This platform also relied on the use of mechanisms developed for detection and correction of outages and not on algorithms optimized for handling fiber delay changes. The vendor of this platform has developed an optimized algorithm that will detect a fiber delay change and very quickly correct all units for it with minimal un-registration. This solution has already been tested by the vendor and is scheduled for implementation in the next system software release.

All platforms were also tested for Requirements 3 and 4 if both switches, FWD and REV, activated simultaneously. The results are listed in Table 4.

Table 4: The Platform Behavior after Optical Redundancy Switch Activation (FWD and REV Simultaneously)

Double Switches		25 km FWD, 25 km REV	50 km FWD, 50 km REV	125 km FWD, 125 km REV	145 km FWD, 145 km REV	Notes
Platform A	Add Delay	3 min 17 sec.	Not tested	Not tested	6 min 45 sec.	Fully loaded HDT modem, no un-reg. of units placed at 0 km
	Reduce Delay	2 min 23 sec.	Not tested	Not tested	7 min 41 sec.	
Platform B	Add Delay	2 min 31 sec.	4 min 8 sec.	3 min 37 sec.	Not tested	No of NIUs was only 8, un-reg. occurred for units at 0 km
	Reduce Delay	2 min 40 sec.	3 min 30 sec.	3 min 23 sec.	Not tested	
Platform C	Add Delay	4 min 13 sec.	4 min 49 sec.	Not tested	Not tested	Fully loaded HDT modem, no un-reg. of units placed at 0 km
	Reduce Delay	4 min 37 sec.	3 min 57 sec.	Not tested	Not tested	

The behavior of the platforms for double protection switching was similar to the behavior when the switchover occurred in one path only. This was expected as the round trip delay change is the parameter of importance. Recovery time for Platform A noticeably exceeded the 4-minute requirement if protection switching added and subtracted 145 km (90 miles) of fiber in both paths. If the switching added or subtracted 25 km of fiber in both paths

simultaneously, the recovery time met the requirement. Platform B performed within the 4-minute recovery requirement, but the tests were again performed with only eight NIUs per modem. As previously, the units at 0 km from HDT (not subjected to switching) un-registered. Platform C also exceeded the 4 minute requirement.

At this time, all three cable telephony platforms do not fully meet

Requirement 3 and Platform B does not meet Requirement 4. All have particular problems in handling certain fiber switch delay differentials. All platform vendors have been encouraged to move from less than optimal detection and correction mechanisms, which were created to deal with specific loss of communication outages, to more robust mechanisms of detection and correction of fiber delay changes. This should result in achieving better recovery results and increased service availability.

Optical Power Changes Accompanying Fiber Protection Switches

During all the tests described, protection switch action was *not* accompanied by a corresponding change in optical power at the input to the receiver. Real networks should be designed to equalize the received optical levels at the AOPS to within ± 1 dB for principal and backup routes. This design consideration is important. If there is an optical amplifier in front of the receiver (after the switch), then power changes will be absorbed by the amplifier and will result only in a slight change in CNR at the receive end. Alternatively, the path losses can be equalized with optical attenuators or couplers. If the level equalization is not achieved, an instantaneous change in RF power from the receiver (and at the input to the HDT modem) will occur. To account for this possibility, all three platforms were tested for the instantaneous level change. All of them showed fast recovery from such changes without affecting service, as long as the changes were within 4-6 dB in RF level (2-3 dB optical change). The call BER was not affected significantly.

This special design requirement for level equalization is not relevant in digital

reverse links (either digitized analog or baseband digital signals). In this case, the changes in optical power levels do not affect receiver RF output levels as long as the optical levels are well within the receive window of the digital receiver.

CONCLUSIONS

New requirements for interactive service platforms with respect to fiber protection switches and increased distances between signal processing equipment (e.g., CMTSS and HDTs) and terminal equipment (e.g., cable modems and NIUs) have been presented and explained. These requirements are partially reflected in DOCSIS standards (absolute and differential distance requirements) but have not been adequately addressed in digital cable telephony platforms. The most current hardware and software releases of the three proprietary platforms deployed in AT&T Broadband markets were tested and analyzed to verify if these requirements were met. The results indicate that all the platforms have problems meeting the recovery requirements after fiber protection switch activation. All of the platforms can accommodate the increased HDT to NIU distance requirements. However, in the test for differential distance requirements, that being the distance between the farthest and closest NIU to the same HDT modem, one of the platforms failed. The vendors anticipate that by the end of this year or in Q1 of 2002 all four requirements will be met. AT&T Broadband will retest these platforms once the upgrades are implemented.

It is critical that similar requirements are extrapolated to DOCSIS-compliant IP platforms, especially the ones supporting voice over IP services, and to PacketCable-compliant platforms.

ADDENDUM – TEST SETUP AND TEST CONDITIONS

To test performance of all digital cable telephony platforms, a full DWDM network was set up in AT&T Broadband

Lab facilities in Westminster, CO, as shown below in Figure 4. Additional testing was performed at vendor locations to simulate more complex situations such as full loading of HDT modems with NIUs.

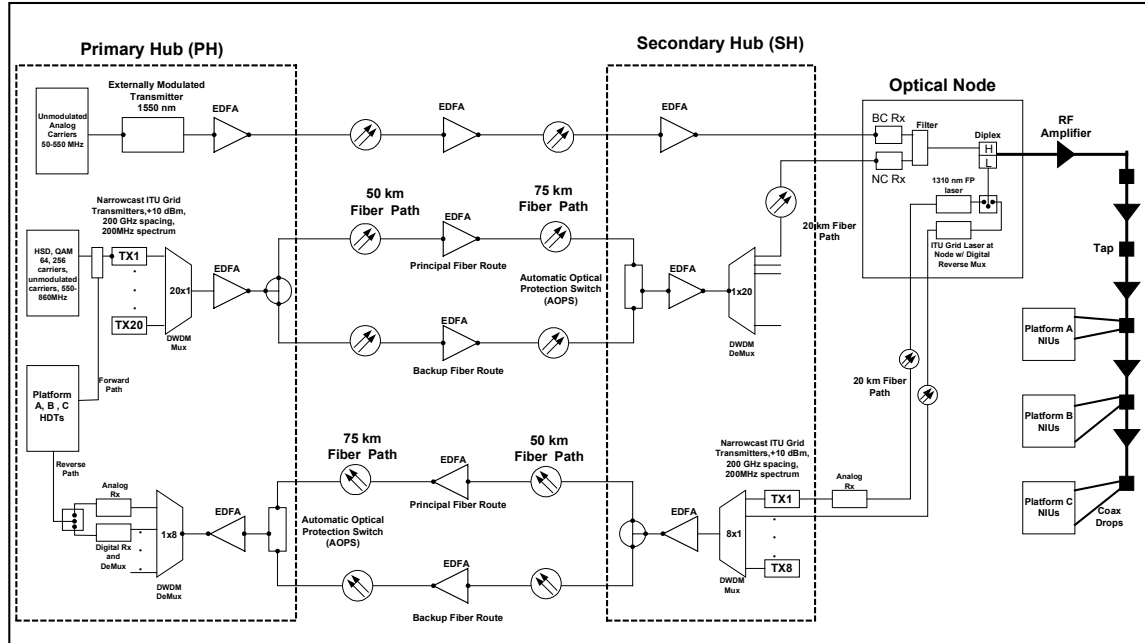


Figure 4: Test Setup in AT&T Broadband Lab

The test architecture consisted of a DWDM system with 20 narrowcast ITU wavelengths (200 GHz spacing) in the FWD path and eight in the REV path. An analog broadcast system carrying 550 MHz of unmodulated carriers on a separate fiber link was also used. The narrowcast signal consisted of 200 MHz of bandwidth carrying the cable telephony signals, as well as some 64 and 256 QAM carriers, HSD signals and several unmodulated analog carriers to fill the band from 550-860MHz. This signal was modulating only one of the narrowcast forward lasers. The narrowcast and broadcast signals were received on separate receivers at a node and combined before RF transport over a simulated coax plant with five amplifiers in cascade. At the end of the cascade, 30 NIUs (for each tested platform) were connected to the network. In the REV

path, both analog and digital return systems were used to carry NIU signals back to their respective HDTs. RF switches were used to select either transport technology in the REV and the platform performance was tested with both technologies. The total distance between HDTs and NIUs was 90 miles (144 km) of real fiber in each direction. This length could be changed to create lower values of total distance by removing fiber spools from the backup fiber routes. The same approach was used to create the differential delays between principal and backup routes. Tests for double (simultaneous) protection switching were also performed. Both AOPSs were activated simultaneously in the FWD and REV paths to reflect real-life situation where the principal fibers for FWD and REV paths are usually in the same fiber sheath and are

usually cut or damaged at the same time. Double switches were also used to attain long round trip delay changes and to simulate situations where switchover results in zero net change in delay.

Details of Test Setup for Requirement 1

Requirement 1 for each platform was tested with NIUs of all types used by AT&T Broadband (e.g., 2-line, 4-line, 12-line, 16-line, locally and network powered and all different hardware types). The NIUs and HDTs were equipped with the currently deployed software releases. The units were connected and powered up in the coax plant at the end of the 90-mile network, one at a time, as they would be under real-life conditions in the field, following the recommended turn-up procedure. Other units were pre-provisioned when possible. The time taken for first-time registration was verified for compliance with the specification. The NIUs were provisioned with dialtone using an Element Manager (EM) or equivalent system interface, and checked for proper operation and RF levels at this maximal distance. They were then subjected to call loads close to the maximal specified calls per hour (cph) for each platform using analog bulk call generators (BCG) to verify call completion rates (CCR) of 99.99% during several overnight/weekend load runs. In parallel, additional tests were performed on individual units to check that all their attributes worked properly at this new distance. Most of the basic and critical features were tested together with those system attributes to verify whether their operation was affected (e.g., turning forward error correction (FEC) on/off, concentration on/off). The tests were performed for many distances to allow for extrapolation of the results over the entire range of distances from 0 to 90 miles.

Details of Test Setup for Requirement 2

To test Requirement 2, several NIUs (of all types) from each platform were connected directly (0 km) to the HDT modems so that their delay from the HDT was negligible, while the other units remained at 70 miles (or 90 miles during testing the objective performance), on the *same* HDT modem. Additional tests were performed with the units closer to the HDTs placed at 600m, 1.1 km, 2.3 km 4.0 km, 25 km, 50 km, 75km, 112 km (70 miles), 125 km and 144 km (90 miles). Load runs were performed with the bulk call generators (BCGs) to test for 99.99% CCR. Provisioning, maintenance and operations functions were performed on the units in parallel to testing for normal operation at 0 miles and 70 miles simultaneously. Further opportunity to test different distances from closest to furthest NIUs from the HDT occurred during testing for Requirement 3 where fiber distance to farthest NIU was changed as a result of the fiber protection switching.

Details of Test Setup and Test Conditions for Requirements 3 and 4

In order to test these two requirements, the lengths of backup paths in the FWD and REV directions were changed by adding and taking away fiber reels in multiples of 25 km. Shorter fiber lengths (1 to 10 km) were used around the critical length change points within each vendor's platform to increase the test resolution. The maximal distance never exceeded the maximal absolute HDT to NIU distance of 144 km (90 miles) of fiber tested before for all three platforms. The optical protection switches were manually activated to cause the determined delay change. HDT modems for platforms A and C were loaded with the maximal number of NIUs. It was not possible to simulate the maximal NIU load for Platform B due to test set up constraints.

This platform was tested with eight NIUs per modem, far below its full load. However, the results were still very meaningful. The NIUs were exercised during the fiber switchovers with bulk call generators at call loads of 6000 to 9000 cph but below the maximal stress specification of the particular platform. Before each protection switch activation, a period of 15-30 minutes of bulk call activity was allowed to establish the CCR. Similarly, a period of 15-30 minutes was allowed to elapse after the switchover to monitor recovery to the previous CCR level and to detect any longer-term anomalies that arose from the switchover. Recovery times were measured from the fiber switch instant to the point at which all units were registered back in the system and the BCGs recovered the CCR levels established prior to the test. The optical power applied to the optical switch from principal or backup fiber paths in FWD and REV was always equalized irrespective of the length differences except when the platform behavior at the instantaneous RF level changes was tested. All platforms had units registered at 0 km on the same HDT modem as units at the longer distances. These were not subject to fiber delay changes.

AUTHOR CONTACT INFORMATION

Ricardo A. Villa, Senior Engineer, DWDM Systems Engineering

Email: ravilla@broadband.att.com

Telephone: (303) 858-5566

Sudhesh Mysore, PhD., Director, DWDM Systems Engineering

Email: smysore@broadband.att.com

Telephone: (303) 858-5204

Oleh Sniezko, Vice President, Engineering

Email: sniezko.oleh@broadband.att.com

Telephone: (303) 858-5027

Acknowledgement

The authors wish to thank Venkatesh Mutalik of Philips Broadband Networks (PBN) for supplying the DWDM equipment used as a network platform for testing. We also wish to thank the three AT&T cable telephony approved vendors for supplying equipment and resources to perform the tests, as well as making their own testing facilities available.

¹ Thomas G. Elliot and Oleh J. Sniezko, "Transmission Technologies in Secondary Hub Rings — SONET versus FDM Once More", *1996 NCTA Technical Papers*: 382.

² Oleh J. Sniezko & Tony E. Werner, "Invisible Hub or End-to-End Transparency", *1998 NCTA Technical Papers*: 247.

³ Tony E. Werner & Oleh Sniezko, "Network Availability, Requirements and Capabilities of HFC Networks", *1996 Conference on Emerging Technologies*: 123.

CABLEHOME™: A HOME NETWORK ARCHITECTURE

Steven Saunders, David Reed, Kevin Luehrs, Stuart Hoggan
Cable Television Laboratories, Inc. (CableLabs®)

Abstract

Home networking is emerging as a highly competitive key strategic market for cable operators. CableHome is a CableLabs initiative focused on enabling cable operators to differentiate by guaranteeing convenient operation of the home network in concert with the delivery of quality services. Creating a supporting infrastructure requires the accommodation of a diversity of networking technologies and protocols, insulation from and integration with non-compliant devices, and a non-restrictive design environment. The CableHome architecture addresses these concerns by facilitating the separation of compliant and non-compliant networks via technology independent logical network elements. These logical elements are correlated with physical devices, non-restrictively, to provide a concrete reference frame. Supplying robust network management and quality of service capabilities in the home is key and these concerns are specifically explored.

INTRODUCTION

Over the next five years, home networking equipment and services have the potential to emerge as perhaps the most strategic market of interest for cable operators since this segment may control how cable operators and consumers can manage the delivery of Internet-based services in the home. Success for cable operators in this market will require a cooperative environment between content producers, service providers, and device manufacturers so that consumers can obtain one-stop shopping and customer care solutions. Indeed, consumers will strongly equate convenience with consumption. In this view, whoever can assemble a home networking system that is easy to deploy and utilize will likely win a large share of this emerging market.

It follows that the cable industry would benefit by offering consumers the following proposition: *If you choose to use home networking equipment approved by your cable operator, the operator will guarantee it will work seamlessly with your broadband cable services delivered over cable.*

Through the CableHome project, the cable industry will establish the technical infrastructure for home networking equipment that will allow individual cable operators to fulfill this promise since today's off-the-shelf products fall short. In particular, with CableHome products consumers will have a guarantee that service quality for their IP telephony, streaming media, or other packet-based services would be maintained over their compliant home networks, and that cable operators would assist them in diagnosing and formulating a response if they encounter problems.

To execute this strategy requires the industry to specify how home networking equipment will interact with the cable operator's "IP cloud". This requires creating specifications that define the minimum set of industry requirements for network management, quality of service, service provisioning, and security functions. Other benefits to cable operators from these specifications should be: 1) lower installation costs by simplifying the home networking installation process; 2) lower equipment costs to consumers through the multiple suppliers enabled by CableHome's open specifications; and 3) lower operating costs by providing cable operators with tools that facilitate remote troubleshooting of consumer problems.

This general set of capabilities allow CableLabs' member companies significant flexibility to develop their own path of

businesses designed to generate revenues from home networking services.

Within this context, several other strategic observations are relevant:

- *Window of Opportunity.* As the broadband data market starts to mature, there still remains a significant window of opportunity for the cable industry to develop a path for customers to obtain convenient home networking products and subscription-based services tied-in with its high-speed data deployments.
- *Threat of Inaction: Cable's a "Bit-Pipe".* Because home networking and the Internet lets any service provider be the "last microprocessor" reaching the customer, thus establishing a dominion that conveniently provisions and manages Internet-based services over home networks, the cable operator's role could be reduced to one of a "bit pipe" or supplying an Ethernet jack.
- *Competition Moving Forward.* Cable's competition is active in formulating their own strategies for providing convenient home networking services and technologies.
- *Open Specification Process.* The CableHome specification process mimics those established by DOCSIS, PacketCable, and OpenCable. This open process allows for any vendors with an interest in developing innovative products for this arena to participate in the development of the specifications, with an eye towards building upon the substantial development already available in the market.

CABLEHOME CHALLENGES AND SOLUTIONS

The home network presents a complex, diverse, and rapidly evolving service delivery environment. Myriad transport technologies, a

variety of device discovery and management protocols, and several structured service delivery platforms all are jockeying for consumer favor in the nascent home-networking marketplace.

The challenge of the CableHome project is to bring some degree of order to this uncertain environment. The end game is to enable the delivery of high-quality cable services to the home, and to design a system to do so requires the consideration of a number of thorny issues.

ARCHITECTURAL CHALLENGES

Choosing a Transport Technology

The first question likely to enter the mind of a network architect when considering the home networking space is: "Which transport technology should be used?" There are many choices: phone line, power-line, new line, and wireless. Each of these choices has strengths and weaknesses, and selecting any particular technology will involve making tradeoffs and concessions.

A tool that has served well in the past in analyzing these types of tradeoffs is the identification of "typical usage scenarios". This turns out to be an elusive exercise because the types and combinations of services consumed on a home network will vary widely across the consumer spectrum. The optimal technology choice is going to vary among cable operators and from consumer to consumer. In the long run, the consumer panacea may well turn out to be a combination of technologies in a single home network. And of course, neither the cable operator nor the consumer wants to get locked long term into a dead-end technology solution.

For these reasons the CableHome project is developing a technical approach that is independent of any particular transport technology. We are able to realize this objective by choosing protocols that operate at the IP layer and above. Insulating quality of service (QoS) mechanisms from the Layer 1/Layer 2 specifics of the various home networking technologies is a particularly challenging task, and is discussed in more detail in a section to follow.

Providing Managed Service Delivery

Two of the primary tenets of the CableHome project are:

- *Operation of the home network and provisioning of services and devices must be extremely simple.*
- *A wide variety of service types with quality guarantees must be enabled.*

In meeting these requirements, a number of complications become apparent. First, “*easy to use*” is relative to technical knowledge, and it is safe to assume that the typical mass-market customer will not possess the expertise nor have the desire to manage the complicated home network environment. With this in mind, a major focus of the CableHome project is the extension of the robust provisioning and network management capabilities of the existing DOCSIS and PacketCable platforms across the home-networking environment.

The second challenge has already been mentioned: there will be many different services with a variety of resource demands used in many different combinations. In order to ensure the quality delivery of these services in numerous and unpredictable usage scenarios, the second major focus of the CableHome project is to extend DOCSIS and PacketCable quality of service mechanisms into the home, providing an end to end QoS solution.

With these two areas of focus in mind, CableHome is drafting network management and QoS specifications (among others) that will define the capabilities needed to ensure the high-quality delivery of a variety of cable services.

We must now consider how to structure the complex home environment such that participating devices can be managed to the

degree necessary to provide quality guarantees. The CableHome approach is to create a clear separation between cable operator-managed networks and non-managed networks in the home. Devices that provide CableHome management and CableHome QoS capabilities are said to reside in the CableHome domain. As shown in Figure 1, CableHome domains are diagrammatically represented as shaded regions and are meant to serve as a visual tool to identify compliant devices that are manageable by cable operators and that are able to take advantage of cable-based service offerings.

CableHome QoS characteristics can be established dynamically. CableHome networks must be insulated from possible performance corruption by non-QoS-capable devices (i.e., devices that reside outside of the CableHome domain). The CableHome architectural element that allows CableHome domains to be established and insulated is known as the **Boundary Point (BP)**, also shown in Figure 1.

The BP is a functional entity (agent) that resides within a device, and it performs the following tasks:

- Supplies the functionality needed to support CableHome QoS and Management;
- May connect non-compliant networks to the CableHome network;
- Insulates CableHome networks from performance corruption that might otherwise be caused by non-compliant devices;
- May serve as a proxy for connected non-compliant devices, enabling indirect management and service delivery to these devices.

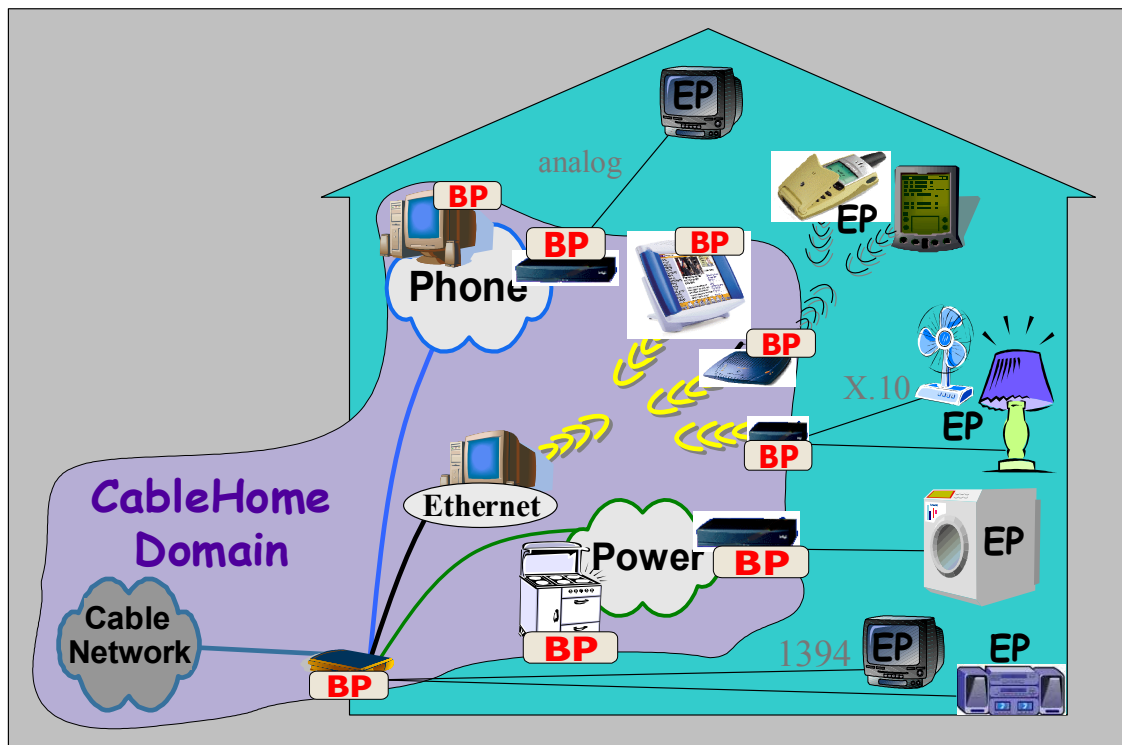


Figure 1: Isolating the CableHome Network

Boundary points generate and consume CableHome messaging and will take the appropriate actions based upon the messages received. BPs always lie on the boundary of the CableHome domain.

The CableHome environment will clearly have to integrate legacy and non-CableHome compliant devices and networks. The ability to deliver services to non-compliant devices may present attractive opportunities to cable operators. But how can these non-compliant devices deliver services when CableHome capabilities are required to do so? CableHome provides this capability through the BP element.

In the CableHome architecture, non-compliant devices that are connected to BPs are known as **Endpoints (EP)**. Endpoints may source or sink data content, but they reside outside of the CableHome domain. These non-compliant entities may range from simple analog audio and video presentation devices to complex non-

compliant networked devices. A BP might connect CableHome networks to the following example types of EPs:

- embedded EP;
- external EP;
- EPs residing on non-compliant networks;
- EP-like applications.

A CableHome boundary point can be thought of as an agent acting on behalf of one or more EPs, enabling them to consume cable services. In the case of a simple embedded analog EP, a BP may do nothing more than convert IP streams to the appropriate format and pass the data on to the EP for consumer presentation. In contrast, a BP may be connected to a functionality-rich EP, in which case the BP and EP might engage heavily in bi-directional communications.

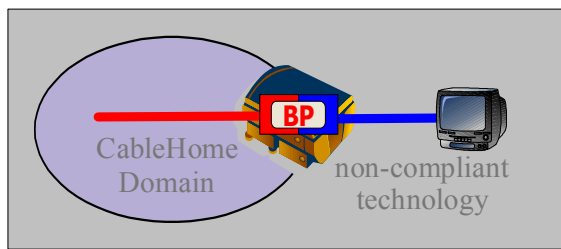


Figure 2: BP Protocol Proxy/Translation

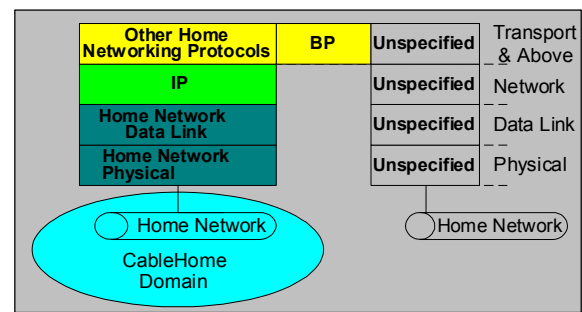


Figure 3: BP Protocol Translation Stack

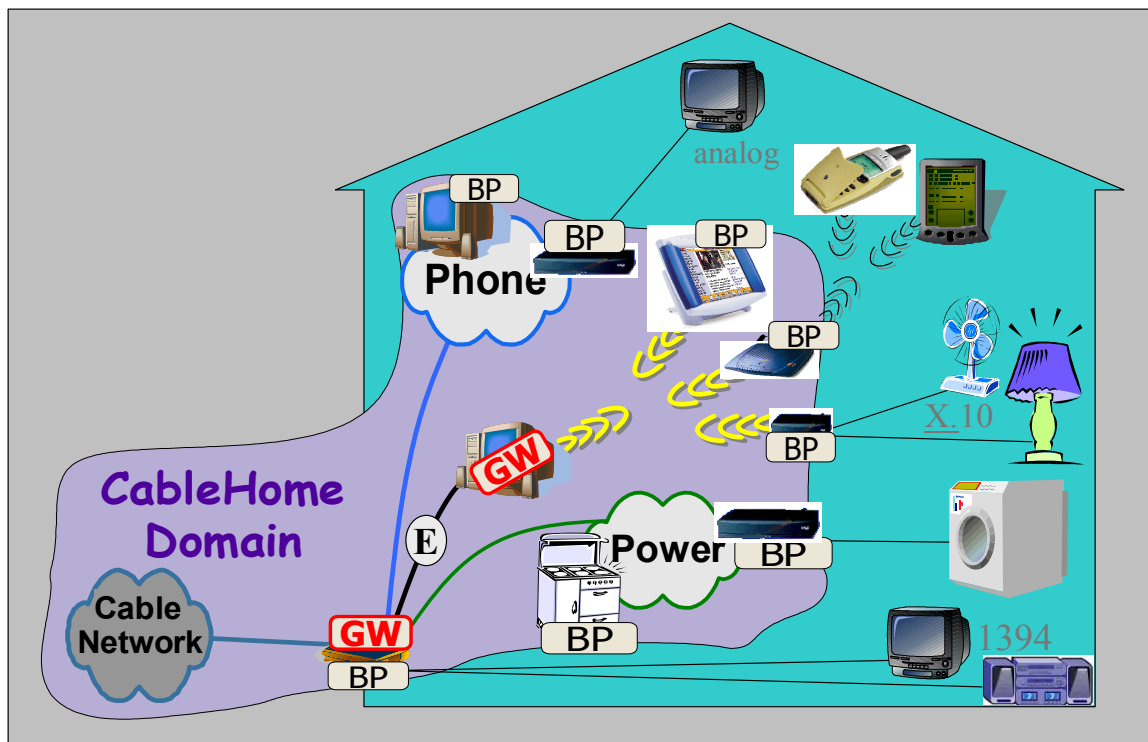


Figure 4: Extending the CableHome Network

BPs may act as a proxy function or as a translation function for the EPs, as shown in Figure 2. The proxy function allows the BP to act on behalf of one or more EPs, while the translation function translates the CableHome-compliant protocols to non-CableHome-compliant protocols. Figure 3 suggests a stack representation for a BP taking on a protocol translation role.

CableHome Traffic Handling

Speculation abounds when it comes to predicting which networking technology will

dominate the home market. Home networking presents a dynamic environment that is ideal for creativity and innovation, and the possible configurations are endless. A reasonable scenario that is occasionally mentioned is a wired “backbone” (for example phone line or power line) with wireless “islands” scattered throughout the house. It clearly is undesirable for a home network architecture to place restrictions on the number, type, or combinations of networking technologies possible in the home.

With all of this in mind, flexibly extending the CableHome domain to multiple home networks becomes an important consideration. The CableHome architectural element that provides this capability is known as the **Gateway (GW)**, shown in red in Figure 4.

Like the BP, the GW is a functional entity (agent) that resides within devices. The GW connects CableHome compliant networks and provides the following capabilities:

- Supplies the functionality needed to extend CableHome QoS and Management capabilities to multiple home networks;
- Insulates the HFC network from in-home traffic;
- Provides routing capabilities required to support unique packet handling needs, such as support for multiple ISP addressing schemes.

Figure 5 demonstrates the placement of a GW propagating CableHome messaging from a wired network onto a wireless network. Figure 6 suggests a stack representation for a GW taking on a such a protocol propagation role.

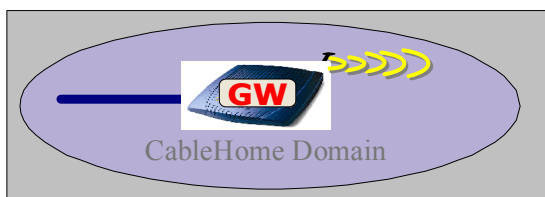


Figure 5: GW CableHome Message Propagation

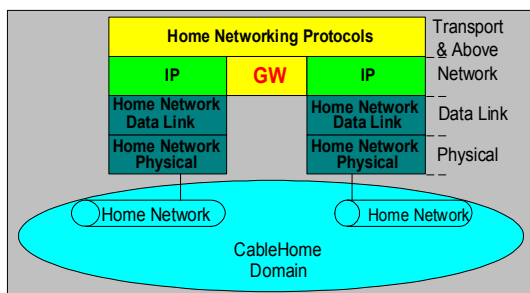


Figure 6: GW Protocol Propagation Stack

Refining the GW and BP Elements

Because the CableHome project is focused both on network management and QoS, specialized versions of the CableHome BP and GW functional elements are identified in the CableHome architecture as follows:

- **MBP: Management Boundary Point**
Provides the previously described BP capabilities for management concerns.
- **QBP: QoS Boundary Point**
Provides the previously described BP capabilities for QoS concerns.
- **MGW: Management Gateway**
Provides the previously described GW capabilities for management concerns.
- **QGW: QoS Gateway**
Provides the previously described GW capabilities for QoS concerns.

In addition, the CableHome domain is further partitioned into separate QoS and Management domains (Q-domain and M-domain respectively). The Q-domain consists of the set of CableHome QoS compliant elements that are able to deliver guaranteed cable-based services. Similarly, the M-domain consists of the set of CableHome management compliant elements that the cable operator can provision and manage.

The Q-domain is defined to be a subset of the M-domain, ensuring that cable operators can manage devices delivering QoS-based services to the degree needed to fulfill service quality guarantees. In addition, the M-domain is defined to extend beyond the Q-domain, allowing CableHome management of products that are not CableHome QoS compliant.

Logical vs. Physical Views

Although not explicitly stated, all of the CableHome architectural elements identified thus far have been **logical** in nature. BPs and GWs are logically bound functional entities that

reside in devices, and do not necessarily imply or specify any specific device or hardware configuration.

The BP and GW logical elements fully define a CableHome network within the home, and they supply all of the in-home functionality defined by the CableHome specifications. Each BP and GW logical element is assigned a unique IP address, and it is the CableHome logical elements that are provisioned and managed. A CableHome network can be conceptualized as a set of BPs and GWs that are discovered and managed, and that interact with each other and with the CableHome support infrastructure, as needed in a flexible manner, to deliver cable-based services.

That said, the CableHome team has found it useful to identify device classes in the home to lend tangible context to the CableHome BP and GW logical elements and combinations of these logical elements. The CableHome concept of device classes places no restrictions on physical devices or combinations of logical elements within physical devices.

There are three classes of CableHome devices, referred to as HA (Home Access), HB (Home Bridge) and HC (Home Client), as shown in Figure 7. The HA, HB, and HC devices are loosely distinguished by their placement in a CableHome network. These device classes provide a strictly informative tool for depicting collections of logical elements but are not considered restrictive. HA, HB, and HC are not addressable entities within the CableHome architecture.

The HA device extends CableHome capabilities from the DOCSIS™ network into the home network. The HA device has a single DOCSIS RF-compliant interface and may have zero or more CableHome compliant interfaces. The HB device extends CableHome capabilities to additional CableHome compliant networks and has at least two CableHome compliant interfaces. The HC device originates and terminates CableHome messaging and has one or more CableHome-compliant network interfaces.

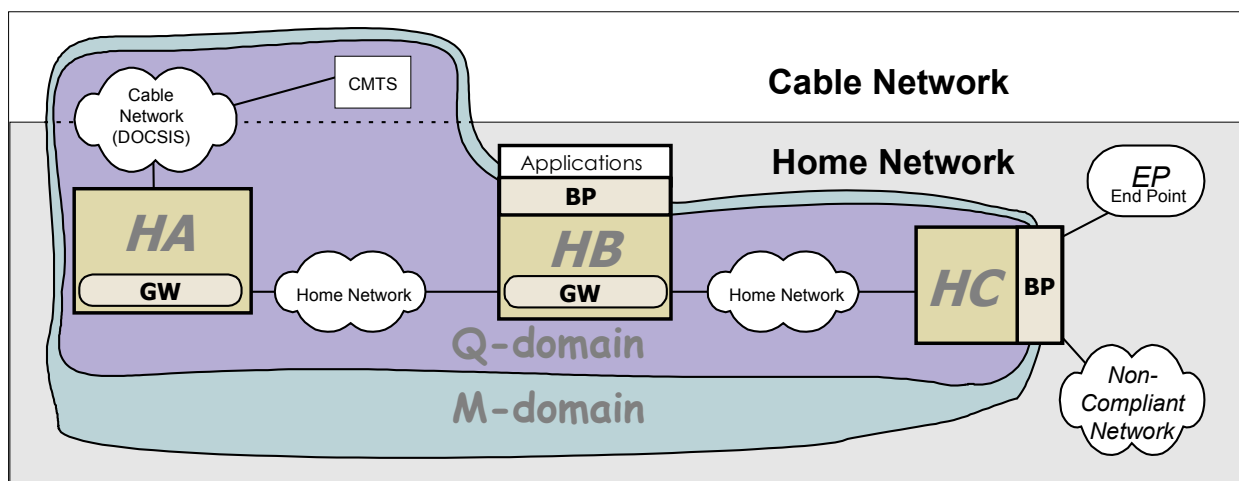


Figure 7: CableHome Device Classes HA, HB, and HC

A CableHome device of a particular type may contain functionality typically associated with other device classes. For example, an HA device, in addition to containing a DOCSIS-compliant interface, may contain HB-like GW

functionality that connects two compliant home networks.

Flexible Solution space

As home networking continues to grow in popularity, consumers may demand many different types of services. To accommodate this demand spectrum, an architecture that allows the design of devices with a broad range of resources and capabilities is critical. As we

have described above, the CableHome architecture provides just such a flexible and non-restrictive design space. Figure 8 demonstrates the flexibility of this approach, showing various example combinations of CableHome elements laid out in various topologies.

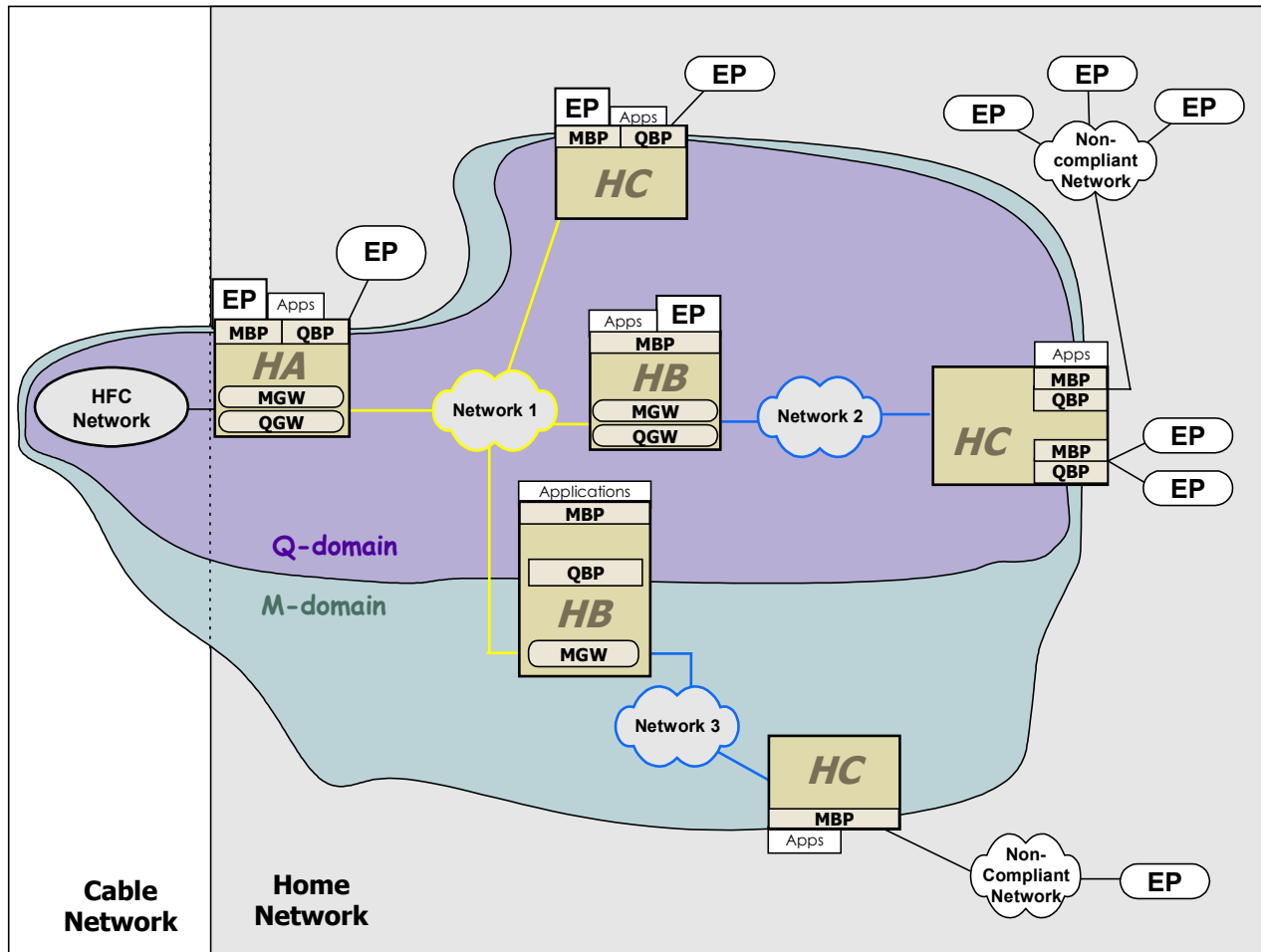


Figure 8: CableHome Architecture, Putting it all Together

CableHome Architecture Summary

The CableHome architecture was developed via a cooperative effort with contributions provided by a large number of service and equipment vendors, in concert with the requirements communicated by cable operators. In summary, the benefits provided are many, including:

- Clearly identified CableHome domain that enables cable operators to manage networks to the degree necessary to deliver quality-guaranteed services.
- A logical element focus that places minimal restrictions on device capability configurations.
- The ability to interface with non-compliant technologies. This enables

service delivery to existing and future proprietary technologies, to non-IP-based network technologies, to thin client devices, etc.

- Interoperable products within the CableHome domains.
- Transport technology independence, leaving a broad set of choices for cable operators and for consumers
- Transport technology independence, allowing cable operators to leverage the existing and rapidly progressing home networking technology base, to deploy new services.
- Flexible implementation space both fosters vendor innovation and permits graceful evolution to include new technologies and services.

MANAGEMENT AND PROVISIONING CHALLENGES

The ability to manage network elements in the home is a critical component of the CableHome vision of guaranteeing reliable delivery of cable-based services in subscribers' homes. The extent to which network elements can be initialized and configured for operation without user intervention is expected to have significant influence on how well CableHome-compliant networks are accepted by consumers. These issues are among those to be addressed by CableHome network management and provisioning systems. This section considers important challenges to the development of CableHome network management and provisioning systems, and how these challenges might be met.

Scope of the Provisioning Task

There may be a great number and type of networked devices in the home that will need to be initialized and configured for operation. Many will likely require unique initialization parameters and configuration processes. The number and type of networked devices in the

home that will need to be initialized and configured for operation is potentially very large. Many will likely require unique initialization parameters and configuration processes. CableHome cannot begin to identify, much less specify, the provisioning needs of all current and future home networking devices and functions.

CableHome is addressing this challenge by limiting consideration of provisioning only to CableHome-compliant devices, and by taking a phased approach to specifying provisioning and management of CableHome-compliant functions. The initial phase will focus on applying the DOCSIS provisioning process to CableHome-compliant devices in order to enable delivery of basic data transport service to these elements connected to the home network. Later phases will explore more complex issues such as service and device capability discovery, and ways to provide basic support specific to a service type such as video distribution.

Management Capabilities

The wide array of devices potentially connected to the home network makes it difficult if not impossible to predict the set of management capabilities needed. "Thin" network devices may have little or no management capability. Some types of devices, such as white goods and other appliances, may have very specialized management needs.

In order to bound the task and limit the scope, the CableHome management system will focus on the management of CableHome GW and BP functions only. The initial set of management functions will be similar to some of the management functions specified in the DOCSIS specifications' suite.

Other network management and provisioning challenges to be addressed during the initial specification phase are listed below:

- Should device and capabilities discovery be supported as part of the CableHome provisioning process?

What protocol(s) should be used and how does it integrate with the existing management protocol framework?

- What do the details of an end-to-end provisioning process for any new service need to be? How are services requested, granted, delivered, billed for, and terminated?
- How will devices on different sub-nets in the home (assigned IP addresses from different service providers) exchange data with one another?
- How will addressing be accomplished as devices are added to the home network, or as leases expire, during periods of isolation from server functions in the cable network (headend or NOC)?
- How will service-level diagnosis be accomplished? How will problems with service establishment, delivery, and termination be detected, reported, and corrected?
- How will CableHome management systems integrate with DOCSIS™, PacketCable™, and OpenCable™ management systems in the cable operator's network?

QUALITY OF SERVICE CHALLENGES

Extending quality of service capabilities from the cable network to the home network presents its own set of unique challenges in order to ensure consistent, reliable service delivery to consumers. Some of these challenges are listed below:

- Maintaining HFC service quality levels as they are delivered over different home networking technologies.
- Developing a common QoS signaling method. This for home network devices and applications to establish prioritized service levels across different, lower-

layer home network transport technologies that interface to DOCSIS 1.1 and PacketCable QoS mechanisms.

- Delivering QoS levels to non-compliant devices and network technologies.
- Establishing a method for mapping DOCSIS 1.1 QoS levels to the QoS levels of different home networking technologies.
- Developing an architecture that is independent of all the different home networking topologies and configurations that are expected.
- Classifying performance characteristics of home networking devices/technologies on their ability to deliver different types of services.

LOOKING FORWARD

This paper offers a description of the vision and challenges incumbent upon establishing a technical infrastructure in the home-networking environment capable of delivering cable-based services in seamless fashion to customers. CableLabs, in cooperation with its member companies, is now working with a large community of vendors to create the technical specifications to implement this vision.

Looking forward, we anticipate new developments in a number of key respects:

- *Entertainment video.* The unique transport requirements for entertainment video applications will likely require modification to existing home networking management and transport protocols. For example, one such immediate requirement is the establishment of a copy protection system for digital media in the home-networking environment.
- *Higher bandwidths.* The increasing bandwidth capabilities of emerging home networking protocols offers cable

operators new opportunities to exploit the broadband capabilities of their networks. Through the CableHome project, the cable industry will work to develop these technologies in a way that will enable the delivery of cable services to consumers.

- *Security.* CableHome will also work to extend onto home networks comparable security mechanisms now provided over cable networks by the DOCSIS and PacketCable specifications.

DESIGN AND PERFORMANCE OF A FULLY-DIGITAL DOCSIS CMTS RECEIVER

F. Buda, E. Lemois, A. Popper, J. Boutros, G. Karam, and H. Sari
Pacific Broadband Communications

ABSTRACT

This paper describes a DOCSIS-compliant cable modem termination system (CMTS) receiver architecture with very advanced features. All receiver functions are implemented digitally, and this feature, together with the advanced signal processing techniques used, leads to an ultra compact and highly scalable CMTS. The receiver architecture adopted makes it possible to implement in a single-chip several upstream burst demodulators along with the corresponding downstream modulators and medium access control (MAC)-layer functions. We report simulation and measurement results confirming the extremely high performance of the described CMTS receiver both in the QPSK and the 16-QAM modes.

1. INTRODUCTION

Hybrid fiber/coax (HFC) networks, which were originally used for broadcast TV services, have recently evolved to two-way networks that deliver high-speed Internet access to residential users. The customer premises equipment in this application is referred to as cable modem (CM) and the network side equipment is called cable modem termination system (CMTS). Potential technologies for high-speed Internet access are asymmetric digital subscriber loops (ADSL) over twisted-pair telephone lines, satellite access, broadband fixed wireless access, and HFC network access. The CM technology has taken the lead among all those technologies, and in the US alone there are today millions of households connected to the Internet over HFC networks.

Standardization for digital data services over HFC networks was undertaken in the past by several organizations including the Digital Video Broadcasting (DVB) project, the Digital Audio-

Visual Council (DAVIC), and the IEEE 802.14 Group. But the slowness of officially accredited standardization groups incited cable operators in the US to form the Multimedia Cable Network Systems (MCNS) consortium in 1995 and define a standard called Data over Cable System Interface Specification (DOCSIS), which has become the *de facto* industry standard in the US. Now, there is also a European version of this standard called Euro-DOCSIS.

The key success factors to any broadband access technology are performance, equipment size, and cost. In the past, there has been a significant effort to integrate and lower the cost of customer premises equipment (CM's), but little effort has been made to reduce the size and cost of CMTS equipment. This is essentially due to the fact that a single CMTS traditionally serves a large number of subscribers, and the cost per user is typically small. This reasoning, which is common to all point-to-multipoint systems, no longer holds when the number of users per network access point becomes small. This is precisely the situation with HFC networks as fiber nodes shrink in size and get closer to the subscribers, and the number of subscribers per port gets smaller. The purpose of this paper is to describe a fully digital CMTS receiver architecture that leads to a very compact and flexible implementation of the CMTS, while ensuring excellent performance.

The paper is organized as follows: First in the next section, we briefly review the DOCSIS Physical (PHY) Layer specification. Next, in Section 3, we present the receiver architecture and describe in some detail the digital front-end and the digital demodulator functions. Section 4 reports simulated and measured performance results of the receiver in the initial ranging mode and in the traffic mode. Finally, we give our conclusions in Section 5.

2. THE DOCSIS STANDARD

DOCSIS is a set of technical specifications [1], [2], which were structured under the leadership of CableLabs to guarantee multi-vendor interoperability. The DOCSIS RF specification includes the physical-layer, the data link control (DLC) layer that includes the medium access control (MAC) sublayer, as well as a convergence layer with upper network layers. Since the topic of this paper concerns the uplink receiver in the CMTS, our description will be limited to the PHY layer with a particular focus on uplink transmission.

2.1. Downstream Channel

Downstream transmission (from the CMTS to CM's) on cable networks can use a channel in a wide spectrum between 50 and 860 MHz. This spectrum is channelized using 6-MHz channel spacing. The modulation format is quadrature amplitude modulation (QAM) with 64 constellation points (64-QAM) or 256 constellation points (256-QAM) [3]. Channel filtering uses a raised cosine filter equally split between transmitter and receiver. The roll-off factor is $\alpha = 0.18$ for 64-QAM modulation and $\alpha = 0.12$ for 256-QAM modulation. The nominal symbol rate on a 6-MHz channel is 5.056941 Mbaud in 64-QAM mode, and 5.360537 Mbaud in 256-QAM mode.

Error correction coding is based on a concatenated coding scheme with an external Reed-Solomon (RS) code [4], [5], an inner pragmatic trellis code [6] and a convolutional interleaver. The RS code used is an RS(128, 122) which is a 3-symbol error correcting code defined over the Galois field GF(128).

2.2. Upstream Channel

Upstream transmission in the DOCSIS standard uses the 5 - 42 MHz frequency band. This spectrum can accommodate a number of upstream channels of different bandwidths. The channel bandwidth W in the DOCSIS specifications can take the values of 200, 400, 800, 1600, and 3200 kHz. The nominal symbol rates for these channel bandwidth values are 160, 320, 640, 1280, and 2560 Kbaud, respectively. That is, the symbol rate is given by $R =$

0.8W. The two modulations specified are the simple quaternary phase-shift keying (QPSK) modulation and 16-QAM. For both modulations, channel filtering uses a raised cosine filter (equally split between transmitter and receiver) with roll-off factor $\alpha = 0.25$. The multiple access scheme is a combination of frequency-division multiple access (FDMA) and time-division multiple access (TDMA), i.e., FDMA/TDMA. In this scheme, the CMTS assigns each CM to one channel and allocates time slots to it on that channel.

The upstream is coded using a Reed-Solomon code over GF(256) with a correction capacity of $T = 1$ to 10 symbols. But there is also an uncoded mode, which corresponds to deactivating this forward error correction (FEC) code. The RS code block length ranges from 18 to 255 bytes, and the number of information bytes per code word ranges from 16 to 253. There are two modes for coding the last block of each burst. The first one is a Fixed Codeword Length which consists of appending by a (0, 0, ..., 0) sequence the last block so that all blocks to the RS coder input are of equal length. The other is Shortened Codeword Length in which the last block is not appended and remains of shorter length than the preceding blocks. The latter mode has the advantage of reducing overhead.

DOCSIS specifications are very flexible in the sense that the modulation format and the FEC code can be defined on a burst-by-burst basis. The burst length itself is redefined at each burst.

2.3. Ranging and Traffic Modes

CM's on HFC networks operate in two different modes: The ranging mode during which different parameters are set, and the traffic mode during which useful data is transmitted. The CM enters the ranging mode at connection set-up (initial ranging) in order to perform carrier synchronization, timing clock synchronization, and power control. In the ranging mode, there is a large uncertainty on these parameters and the search for the optimum parameters must be therefore performed over an extended range. In the traffic mode, the CMTS has some *a priori* knowledge of these parameters, and the uncertainty is small. Synchronization problems must therefore be examined in the ranging mode.

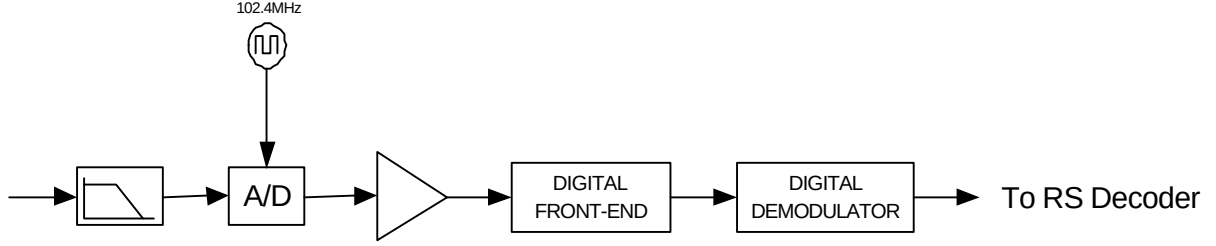


Fig. 1: General block diagram of the CMTS receiver.

Another important function of the CMTS in the ranging mode is to estimate the channel impulse response and compute the optimum equalizer coefficients for that channel. The computed coefficients are then sent to the CM, which is in charge of pre-equalizing the transmitted signal in the traffic mode.

3. RECEIVER ARCHITECTURE

A general block diagram of the receiver is shown in Fig. 1. The received signal is first filtered, amplified, and A/D converted using a clock generated by a free-running oscillator. The nominal frequency of this clock is 102.4 MHz. The variable-gain amplifier used controls the signal power of the entire carrier multiplex. After A/D conversion, the signal is sent to the fully digital front-end which is followed by the digital demodulator.

3.1. Front-End

A functional block diagram of the digital front-end is shown in Fig. 2. The first function of this block is to convert the received digital signal to baseband and generate the in-phase (I) and quadrature (Q) baseband components. This is performed using two multipliers and a numerically controlled oscillator (NCO). The frequency of this oscillator is controlled by the CMTS so as to extract the desired carrier. This signal is then passed to digital filtering and decimation stages, which provide 4 samples per nominal symbol duration. The final stage of the digital front-end is the matched filter, which operates at 4 times the nominal symbol rate and performs square-root raised-cosine Nyquist filtering.

3.2. Digital Demodulator

The front-end is followed by the digital demodulator whose basic function is to perform timing and carrier synchronizations, channel equalization, ingress noise cancellation, and make symbol decisions. A functional block diagram of the demodulator is depicted in Fig. 3.

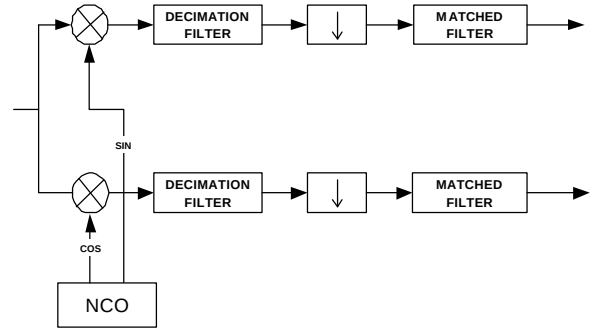


Fig. 2: Block diagram of the digital front-end.

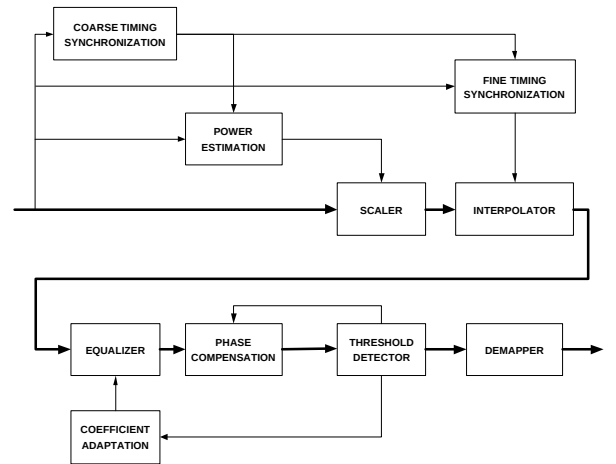


Fig. 3: Block diagram of the digital demodulator.

3.2.1. Coarse Timing Estimation

First, a coarse timing function detects the beginning of each burst with the required precision (typically a precision of half a symbol period). The conventional approach to coarse timing estimation is based on power estimation. The principle is simple: The signal received in the absence of bursts is due to noise, and its value is small compared to the signal received during bursts. Therefore, a power estimation circuit followed by a threshold comparator gives an indication of the start of bursts. The power estimation circuit is composed of two elementary functions: The first one is a squaring circuit which gives the instantaneous signal power, and the second is a low-pass filter which performs short-term averaging.

The first problem associated to this concept is that the precision of the burst start estimate is a function of the filter used. A short filter memory is required to improve precision, but then the estimator becomes very sensitive to additive noise. That is, robustness to noise and precision of the estimator are two contradictory requirements in this technique. The second problem is that the threshold is a function of the received signal power level, which is undesirable. A high threshold leads to the risk of missing bursts and includes an estimation delay. A low threshold reduces the delay, but creates the risk of declaring that a burst is present when no burst is actually transmitted. (This occurs when the noise power exceeds the threshold level.)

To avoid these problems, we developed a new coarse timing detector that involves a correlator and the computation of a *contrast function* that is independent of the received signal power level. The correlator correlates the incoming signal with the preamble sequence stored in the receiver. (The preamble must have good correlation properties, i.e., a very narrow correlation peak and very low correlation values around that peak.)

With a *contrast function* that is independent of the received signal power level, a fixed threshold can be used (without any performance penalty) to detect the correlation peak and the burst start. The threshold comparator in the block diagram determines a short time-window in which the

correlation maximum is to be searched. In the traffic mode, the CMTS has some *a priori* knowledge of the burst position and knows the time window over which the *contrast function* needs to be maximized.

3.2.2 Subsequent Demodulator Functions

Next, a fine timing function determines the right sampling instant and passes this information to an interpolator that generates symbol-spaced signal samples. A scaler that precedes the interpolator sets the power of the over-sampled signal to a predetermined value. The scaler is controlled by a power estimation function that is activated during signal bursts. The symbol-spaced signal samples generated by the interpolator are passed to subsequent receiver stages, which include an adaptive equalizer, an ingress noise canceller, and a carrier phase recovery circuit.

The equalizer is a linear equalizer whose coefficient values are computed using the zero-forcing (ZF) criterion [3]. This criterion is more appropriate in the present case than the more popular minimum mean-square error (MMSE) criterion, due to the requirement to send the coefficient values to the CM to implement a pre-equalizer. The reason is that the MMSE equalizer makes a trade-off between channel distortion and additive noise, and therefore the computed coefficients do not perfectly invert the channel transfer function. And contrary to an equalizer at the receiver, a pre-equalizer does not amplify the additive noise. Therefore, the best coefficient setting for the pre-equalizer is that which perfectly inverts the channel.

As is well known, ingress noise represents one of the major disturbances that affect upstream data transmission in HFC networks. Ingress noise is essentially due to local AM radio signals and other types of disturbances that leak into the cable. It is modeled as narrowband interference that may be on or off and essentially constant over a period of time which can be in excess of several minutes. Another characteristic of ingress noise is that, contrary to channel distortion, which is specific to each CM, it is common to all CM's sharing the same upstream carrier. The reason is that the CMTS receives the

sum of all noises that leak into the cable at all customer premises that it serves, and the resulting noise equally affects all time slots no matter where they originate from.

For reliable data transmission on the upstream channel, the receiver must include an efficient ingress noise canceller, particularly for 16-QAM and higher-level modulations. One way to suppress ingress noise is to use a notch filter at the ingress noise frequency, but notch filtering also distorts the useful signal and creates intersymbol interference (ISI), which is undesirable. An alternative approach consists of estimating ingress noise by means of a prediction filter and subtracting this estimate from the received signal prior to threshold detection. The latter approach, which leads to significantly better performance, was adopted in our receiver design.

The final function before the threshold detector (which makes the symbol decisions) is the carrier synchronization function. This includes a frequency estimator that estimates the frequency offset between the CM and the CMTS and a phase recovery circuit that synchronizes the carrier phase of the incoming signal. The estimated frequency offset is used to derive a control signal that is sent to the CM to synchronize its oscillator frequency with that of the CMTS. The decision-feedback frequency estimator used in our design is based on a newly developed algorithm that is very robust against symbol decision errors. Finally, the phase recovery circuit compensates for residual synchronization errors between the CM and the CMTS.

4. PERFORMANCE RESULTS

Performance of the designed CMTS receiver was evaluated using extensive computer simulations and laboratory measurements. In this section, we will report results that assess the performance of the digital front-end and of individual demodulator functions, as well as results that assess overall receiver performance.

4.1. Front-End Performance

Performance of the CMTS receiver was tested for different symbol rates and different loads of the upstream spectrum. The most unfavorable condition

for the digital front-end occurs when the desired signal has the lowest symbol rate (160 kbaud) and arrives at the receiver with the minimum signal level, while adjacent carriers arrive with the maximum signal level allowed in DOCSIS specifications.

To evaluate worst-case performance, we have simulated a carrier multiplex where a 160 kbaud desired signal is received together with 8 adjacent carriers each having a data rate of 2560 kbit/s and a power spectral density (psd) that is 12 dB above that of the desired signal. That is, the desired signal power was 24 dB below that of each one of the other carriers. Our simulations have indicated that the receiver performance in these conditions is essentially the same as in the case of an isolated carrier over the upstream channel. The measurement results using a lab prototype were very much in agreement with the simulation results.

4.2. Overall Performance in the Traffic Mode

Next, we simulated the overall bit error rate performance (BER) of the receiver in the traffic mode and plotted it as a function of the signal-to-noise ratio (SNR). The results take into account the total imperfections of the receiver including those of the front-end and of the synchronization functions. The BER vs. E_b/N_0 (transmitted energy per bit to the noise spectral density ratio) curves are given in Fig.4 for QPSK and in Fig. 5 for 16-QAM. The figures also show the theoretical BER curves which correspond to the performance of an ideal modem.

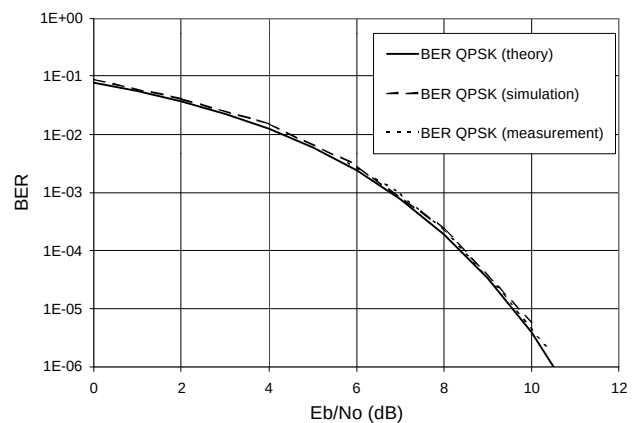


Fig. 4: Overall BER performance of the CMTS receiver in the QPSK mode.

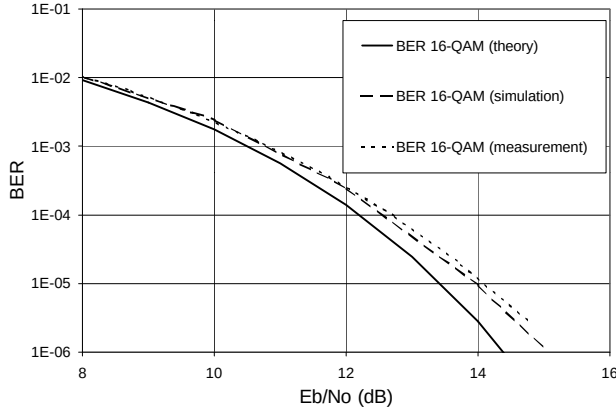


Fig. 5: Overall BER performance of the CMTS receiver in the 16-QAM mode.

These figures show that the overall degradation at the BER of 10^{-6} is limited to 0.2 dB in QPSK and 0.6 dB in 16-QAM. These results are obtained in the absence of error-correction coding. This means that the 0.2 dB SNR degradation in QPSK and 0.6 dB degradation in 16-QAM will hold for BER values as low as 10^{-10} or 10^{-12} after RS decoding.

Figs. 4 and 5 also give the measurement results. We can see that the measurement results coincide with the simulation results in QPSK, and that the difference between simulated and measured results is limited to 0.1 dB in 16-QAM.

4.3. Burst Detection in the Ranging Mode

The most significant performance indicator in the ranging mode is the time needed by a CM to register with the network. The registration time must be evaluated in two extreme cases: The worst case, which corresponds to all CM's using the ranging opportunities, a situation which typically occurs after a CMTS Reset, and the best case, which corresponds to only one CM using the ranging opportunity. The latter case is in fact sufficient to determine the performance of the CMTS receiver.

The important parameters to consider here are the elementary probabilities Pnd_{el} and Pfa_{el} which respectively correspond to missing a burst and to a false alarm at a given time t . Missing a burst occurs when the burst is actually transmitted and the *contrast function* used for burst detection takes a

value lower than the decision threshold S . A false alarm corresponds to the contrast function taking a value that exceeds the decision threshold while no burst is actually transmitted. Note that Pfa_{el} is independent of the SNR, because in the absence of useful signal, both the numerator and the denominator of the contrast function $C(t)$ are proportional to the noise variance, and therefore the noise variance cancels out. In contrast, Pnd_{el} is a function of the SNR. To reduce Pnd_{el} , we need to decrease the threshold S , and to reduce Pfa_{el} , we need to increase S . That is, reducing the elementary non-detection probability is a contradictory requirement with reducing the elementary false alarm probability. This is shown in Fig. 6 where we have plotted Pnd_{el} for $E_b/N_0 = 10$ dB and $E_b/N_0 = 8$ dB.

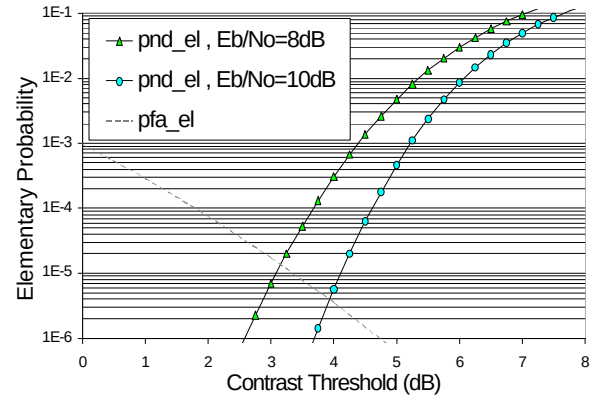


Fig. 6: Elementary false alarm and non-detection probabilities.

But Fig. 6 alone is not sufficient to determine the optimum value of the threshold S . To determine this value, we need to consider the full probability of missing a transmitted burst, which we denote Pnd . To evaluate this probability, we need to consider the following two situations:

- The transmitted burst is not detected due the *contrast function* taking a value lower than the threshold. This occurs with a probability of Pnd_{el} .
- A false alarm occurs during one of the N symbols preceding the burst start, where N is the number of symbols in the ranging burst. A false alarm will activate the demodulator and deactivate the *contrast function* calculations for the following N symbols, and therefore a true

burst start during N symbols after a false alarm will not be detected. The probability of missing a burst due to false alarms is therefore $N.Pfa_{el}$.

Taking into account, these two types of missing a burst, the total probability of missing a burst is given by

$$P_{nd} = P_{nd_{el}} + N.Pfa_{el}.$$

According to DOCSIS specifications, a ranging burst must carry at least 34 bytes, which map on 136 QPSK symbols. Taking into account this minimum value as well as the preamble and the redundancy for error-correction coding, the number of symbols in the ranging burst is close to 200. Therefore, $P_{nd} = P_{nd_{el}} + 200Pfa_{el}$ is a good approximation for the probability of missing a transmitted burst, and the threshold S must be set so as to minimize this probability. In Fig. 7, we have plotted the probability of missing a burst as a function of the threshold value, for $E_b/N_0 = 10$ dB and $E_b/N_0 = 8$ dB.

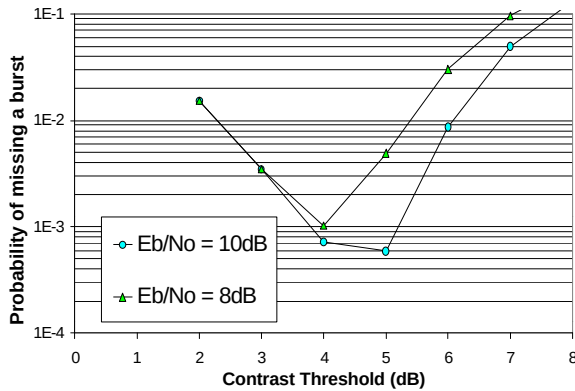


Fig. 7: P_{nd} as a function of the threshold.

This figure indicates that any value of the threshold between 3 and 4 will give a P_{nd} lower than 10^{-2} for E_b/N_0 values higher than 10 dB. A P_{nd} lower than 10^{-2} means that the CM will use the first ranging opportunity 99% of the time, and will need at least two ranging opportunities with a probability of 10^{-2} . Continuing further, the CM will need 3 ranging opportunities only with a probability of 10^{-4} . Consequently, the average registration time of a CM (assuming that only one CM is trying to register at a time) is well approximated by 1.01

times the interval between two consecutive ranging opportunities.

4.4. Pre-equalizer Performance

As mentioned earlier, the pre-equalizer coefficients are computed by the CMTS receiver and sent to the CM to pre-equalize the transmitted signal in the traffic mode. Here, we give the performance corresponding to two different options. In the first, the coefficients are estimated using a single ranging burst, and are immediately sent to the transmitter after that burst. In the second, two ranging burst are used in order to obtain a better estimate of the optimum pre-equalizer coefficients.

To assess the BER performance of the pre-equalizer, we used a channel model with 3 echoes: The first echo is 10 dB below the main signal path and has a delay of 1 symbol period T , the second echo is 20 dB down and has a delay of $2T$, and finally, the third echo is 30 dB down and has a delay of $3T$. Note that the performance results obtained with this model are independent of the symbol rate. Fig. 8 shows the BER curves corresponding to 16-QAM and an 8-tap pre-equalizer. The figure also shows the theoretical 16-QAM curve corresponding to an ideal modem operating on a channel with no distortion.

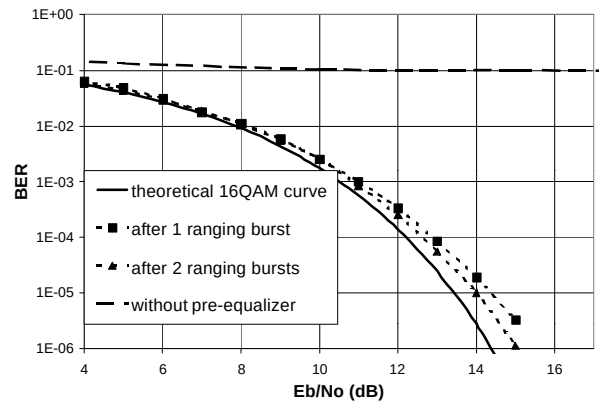


Fig. 8: BER performance of an 8-tap pre-equalizer with 16-QAM modulation.

Notice that without a pre-equalizer, the system has an irreducible BER on the order of 10^{-1} . The results show that the SNR degradation at the BER of 10^{-5} is approximately 1 dB when only a single

ranging burst is used to optimize the pre-equalizer coefficients. This degradation is reduced to 0.6 dB when two ranging bursts are used for coefficient optimization.

Next, we investigated the pre-equalizer signature using the 1-echo channel model which appears in DOCSIS specifications. The signature gives a plot of the echo amplitude (in the dB scale) vs. echo delay (normalized by the symbol period T) that leads to an SNR degradation of 0.8 dB at the BER of 10^{-5} . The results corresponding to 16-QAM and an 8-tap pre-equalizer are depicted in Fig. 9.

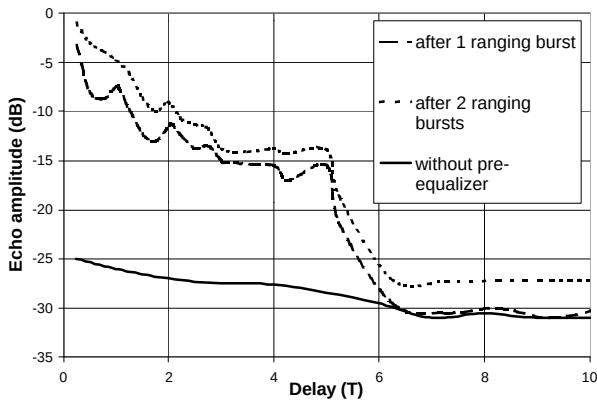


Fig. 9: Signature of an 8-tap pre-equalizer with 16-QAM modulation.

The figure shows that non-equalized 16-QAM system only tolerates an echo of -25 dB (the normalization is with respect to the main signal path). It also shows that for delays up to 5 symbol periods, an 8-tap pre-equalizer can cope with an echo amplitude of -15 dB. The sharp drop of performance is due to pre-equalizer size and reference-tap position used in these simulations. Another interesting observation that can be made here is that using two ranging bursts to optimize the pre-equalizer improves the signature by 1 to 3 dB.

4.5. Ingress Noise Canceller Performance

Performance of the implemented ingress noise canceller was simulated using a 16-QAM signal at the symbol rate of 2.56 Mbaud. The ingress noise model used includes three interferers respectively centered at a distance of 60 kHz, 500 kHz, and 1000 kHz, from the carrier frequency. Each

interferer has a width of 20 kHz and an individual power that is 15 dB below the useful signal power.

The results obtained using this model are depicted in Fig. 10. Notice that a 16-QAM modem cannot operate in the presence of this type of ingress noise without an interference canceller. Clearly, the BER curve shows a floor close to 10^{-1} , which is entirely unacceptable. With a simple canceller based on noise prediction, the BER curve becomes parallel to the ideal 16-QAM curve, and the SNR degradation becomes less than 2 dB. Better performance can be achieved by increasing the number of taps of the noise prediction filter.

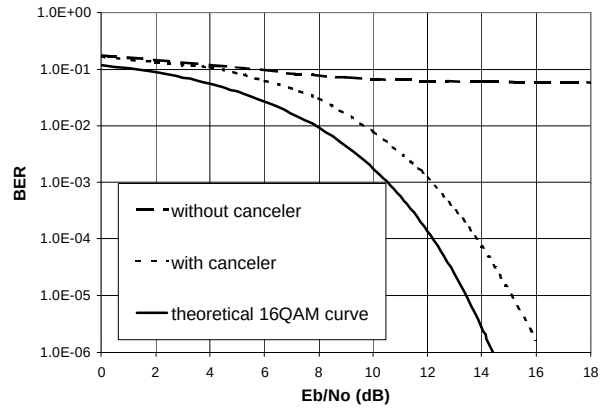


Fig. 10: Influence of ingress noise on 16-QAM and its compensation using noise prediction.

5. SUMMARY AND CONCLUSIONS

We have presented a fully digital receiver architecture, which substantially reduces the size and cost of the CMTS while ensuring excellent overall performance even under worst-case conditions. First, the simulation and measurement results confirmed that the multi-channel front-end gives quasi-ideal performance in the very unfavorable condition where in addition to the useful signal, the cable carries a multiplex of 8 modulated carriers with an individual power that is 24 dB above that of the useful signal. Next, the total SNR degradation of the receiver at the (uncoded) BER of 10^{-6} was found to be limited to 0.2 dB in the QPSK mode and 0.6 dB in the 16-QAM mode. It was also found that with the synchronization algorithms implemented, the average time needed by a CM to register with the network (assuming only one CM is

trying at a time) is only 1.01 times the ranging opportunity, which is an extremely short registration time. Finally, the receiver also includes an efficient adaptive equalizer to compensate for channel distortion and an ingress noise canceller.

The presented receiver was extensively tested using an FPGA implementation and integrated in an ASIC using high-speed 0.18 μ CMOS technology. Prototypes of this chip, which also includes other physical-layer functions as well as MAC-layer functions, are due from foundry in June.

REFERENCES

- [1] Data-over-Cable Service Interface Specifications – Radio Frequency Interface Specification, SP-RFI-105-991105, Cable Television Laboratories, November 1999, Louisville, Colorado.
- [2] Digital Multiprogram Systems for Television Sound and Data Services for Cable Distribution, ITU-T Recommendation J.83 (Annex B), April 1997, ITU, Geneva, Switzerland.
- [3] J. Proakis, “Digital Communications”, Third Edition, McGraw Hill, 1995.
- [4] R. Blahut, “Error Control Codes,” Addison Wesley, Reading, MA, 1984.
- [5] G. C. Clark, Jr., and J. B. Cain, “Error Correction Coding for Digital Communications,” Plenum Press, New York, 1981.
- [6] J. K. Wolf and E. Zehavi, “Pragmatic Trellis Codes Utilizing Punctured Convolutional Codes”, IEEE Communications Magazine, no. 32, pp. 94 – 99, February 1995.

DETECTION AND CLASSIFICATION OF RF IMPAIRMENTS FOR HIGHER CAPACITY UPSTREAMS USING ADVANCED TDMA

Daniel Howard
Broadcom Corporation

Abstract

Mitigation of ingress, common path distortion (CPD) and impulse/burst noise is of great interest to cable operators. Advanced physical layer technologies such as frequency agility, enhanced error correction, ingress cancellation and FFT processing have created new tools for mitigation of impairments. Impairments can now be canceled, reduced or avoided by intelligent adaptation of both signaling and system parameters. A detailed knowledge of RF impairments is required however, to both characterize the impairment and to determine the best mitigation strategy.

In this paper, detailed models of common path distortion (CPD), ingress, and impulse/burst noise are presented based on first principles and verified by plant measurements. The models were developed to optimize the detection, classification, and mitigation and/or avoidance of such impairments by advanced DOCSIS TDMA physical layer technology [1]. The result is higher upstream performance and capacity. A link between plant monitoring and communication system configuration is made that can significantly improve availability and overall plant maintenance. The approach is applicable to existing as well as deep fiber architectures.

INTRODUCTION

Advanced TDMA improves upstream capacity and performance via greater bits/Hz, more robust signaling, and higher signaling bandwidth. In clean upstream channels, a threefold increase in capacity is provided by operation at 64 QAM at 5.12 Megasymbols per second (Msps). In impaired channels, advanced TDMA provides cancellation of ingress and common path distortion (CPD), mitigation of high rate and long duration impulses by greater FEC ($T=16$) and Reed Solomon byte interleaving. Coupling these advanced features with upstream monitoring in both the

spectral and time domains creates a new capability for operators to simultaneously report problems and to adapt the signaling to the impairment in order to maintain the capacity of the system as high as possible.

Previous measurements have concentrated on long term statistical characterization, or capture of short time traces of the upstream only. Further, many of the models currently available are still basic in nature. For example CPD is usually considered to produce beats or at most triplets every 6 MHz. Many more frequencies exist in the CPD spectrum however, and proper design of cancellation filters requires knowledge of this additional spectral structure. Hence, in this paper, more detailed models of common RF impairments are provided in terms of parameters needed by advanced TDMA systems to obtain the maximum capacity from the HFC system. The new models are verified by plant measurements, and then applied to an advanced TDMA system to show how the system may be adapted to the impairments.

MODELING RF IMPAIRMENTS

Common Path Distortion

Common path distortion (CPD), also termed common path intermodulation distortion (CPID) arises in cable plants from several mechanisms, but the most common source is oxidation of contacts which leads to diode-like behavior. A second mechanism is improper balancing of actives leading to nonlinear behavior. In general, the impairment produces second and third order nonlinearities which lead to mixing products of downstream frequencies. These products are located in the upstream frequency band. The most common of these products are at difference frequencies of the downstream video carriers: 6 MHz, 12 MHz, and so on up to 42 MHz. However a complete characterization of CPD reveals many more frequencies in the upstream spectrum,

and also the different bandwidths of these frequencies.

The mechanical discussion of common path distortion (CPD) has been covered texts such as Ciciora et al [2] and will not be repeated here. We begin by assuming that second and higher order mixing products have been produced in the cable plant, and seek to derive the frequencies, relative amplitudes, and fine structure of CPD in the upstream frequency band. It will be shown that there are three scales of CPD spectral structure:

Coarse

Main CPD frequencies that depend on whether the plant is set up for harmonically related carriers (HRC), incrementally related carriers (IRC) or standard carriers (STD). These include the well-known 6 MHz beats in the upstream spectrum.

Medium

Sidebands around each coarse CPD frequency that result from the use of offset carriers in certain cable channels as per FCC regulations for avoiding aeronautical radio communications. These offsets are either 12.5 kHz or 25 kHz away from the nominal downstream frequencies.

Fine

Spreading of CPD coarse and medium frequencies with occasional tone-like peaks that result from carrier frequency inaccuracy in downstream modulators. Typical carrier frequency accuracy of cable modulators is on the order of +/- 5 kHz to +/- 8 kHz.

An NTSC downstream signal has two main peaks, at the video and audio carriers. If f_v is the frequency of the video carrier, then the audio carrier will be at $f_a = f_v + 4.5$ MHz. Generally speaking, subsequent carriers for other downstream cable channels will be at $f_v + m*6$ MHz, $f_a + m*6$ MHz, where $m=1,2,3$, and so on.

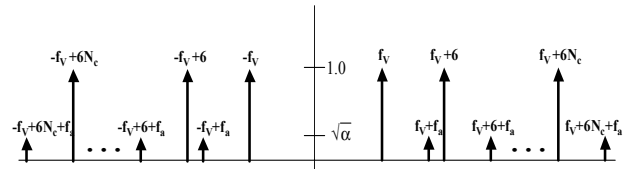
To determine the CPD frequencies that result from second order mixing products, we need only take sum and difference frequencies, $f_j - f_i$, where both positive and negative frequencies of the original spectrum must be considered. The result is CPD beat frequencies at 6, 12, 18... $m*6$ MHz,

with sidebands at +/-1.5 MHz around every 6 MHz beat. Thus, the main, or coarse CPD frequencies from second order mixing products in the upstream band are at 6.0, 7.5, 10.5, 12.0, 13.5, 16.5, 18.0, 19.5, 22.5, 24.0, 25.5, 28.5, 30.0, 31.5, 34.5, 36.0, 37.5, 40.5, and 42.0 MHz. Since these frequencies are invariant to downstream carrier shifts due to frequency plan, they will always be present when CPD exists.

In order to determine relative amplitudes of the CPD frequencies, and for greater detail in modeling, a formalism for the above will be used which is based on the fact that multiplication in the time domain is equivalent to convolution in the frequency domain. Since the frequency domain representation of a real carrier at f_v is $\frac{1}{2} [\delta(f+f_v) + \delta(f-f_v)]$, where δ is the Dirac delta function, if we represent the entire cable downstream spectrum as only the video and audio carriers, the spectrum can be written as

$$S(f) = \sum_{\substack{n=-N_c \\ n \neq 0}}^{N_c} \{ \delta(f-nf_c) + \alpha \delta(f-[nf_c+f_a]) \}$$

where N_c is the number of downstream cable channels, f_c is the spacing between channels (6 MHz), α is the amplitude of the audio carrier relative to the video carrier (-8.5 dB) and f_a is the spacing between the audio carrier and the video carrier (4.5 MHz). This spectrum is depicted in the figure below:



Simplified Downstream Spectrum Model

The second order mixing products can then be determined from

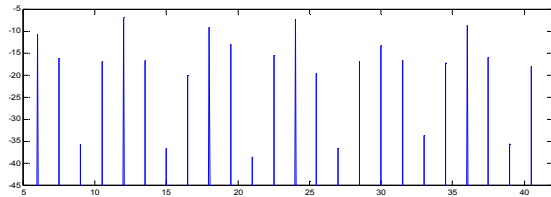
$$S_2(f) = S(f)*S(f)$$

where * denotes convolution. A similar approach is used to derive the 3rd order mixing products:

$$S_3(f) = S_2(f)*S(f) = S(f) * S(f) * S(f)$$

Additional CPD frequencies are produced, for example at $f_k+f_j-f_i$, and also at $2f_j-f_i$ and f_j-2f_i . For HRC systems, these additional frequencies are at multiples of 1.5 MHz since the original carriers are at multiples of 6 MHz + (0 or 4.5 MHz).

A simple MATLAB routine to calculate the up-stream band of CPD 2nd and 3rd order frequencies was used to predict CPD for an HRC system:

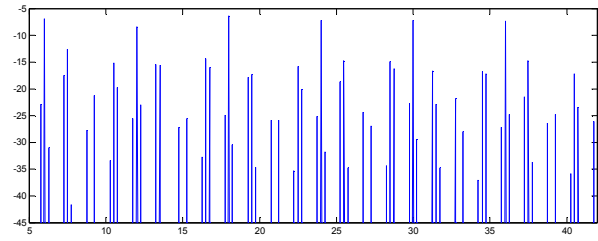


Coarse CPD Frequencies for an HRC plant

The CPD frequencies at 9, 15, 21, 27, 33, and 39 MHz are solely due to 3rd order products, while the remaining frequencies are due to both 2nd and 3rd order products.

Note that STD and IRC plans have carrier frequencies which are offset by 0.25 MHz from those of HRC plans. While this does not affect the location of the 2nd order mixing products, it will affect the location of 3rd order products. For example, in an IRC or Standard plant, the audio carrier of Chanel 19 will be at $151.25 + 4.5 = 155.75$ MHz. Twice the video carrier of Channel 4 is $2*67.25 = 134.5$ MHz. The difference between the two is 21.25 MHz. Hence, a key visual cue for whether the plant is HRC or Standard/IRC is the presence of CPD frequencies at x.25 MHz or X.75 MHz locations; only Standard and IRC plans will produce these coarse CPD frequencies. It is possible that on an IRC plant, the 3rd order frequencies can be much higher due to more coherent summing of the mixing products, however this has yet to be verified with measurements [5].

The MATLAB routine was rerun with IRC frequencies, with the resulting spectrum shown at the top of the next column:



Coarse CPD Frequencies from IRC Plant.

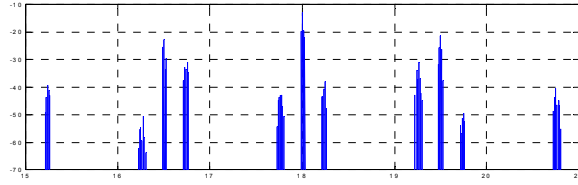
Derivation of Medium CPD Frequency Structure

There are additional medium scale CPD frequencies which are produced by the fact that the FCC requires cable operators to offset the carriers in certain bands by either 25 kHz or 12.5 kHz to prevent any leakage signals from interfering with aeronautical radio communications in those bands. The rules are as follows [3]:

Cable in the aero radiocom bands 118-137, 225-328.6 and 335.4-400 MHz must be offset by 12.5 kHz.

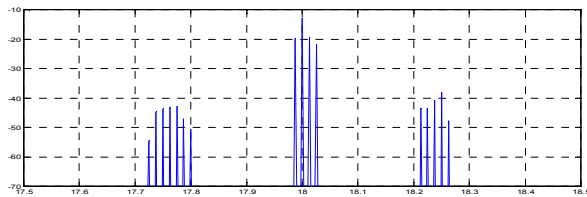
Cable channels in the aero radiocom bands 108-118 and 328.6-335.4 MHz must be offset by 25 kHz.

Second order difference frequencies between an offset carrier and a non offset carrier will thus produce CPD frequencies at 12.5 and 25 kHz offsets from the previously predicted frequencies. Third order offset products will produce additional CPD frequencies at 37.5 kHz, 50 kHz, 62.5 kHz, etc. from the non offset products, which will be lower in amplitude due to the fact that the number of cable channels which must be offset is less than the number which are not offset. The MATLAB routine was rerun for IRC frequency plans using positive offsets for the FCC regulated cable channels with the following result:



Medium Scale CPD Structure on STD Plant

The FCC offsets result in widening the CPD tones via additional tones from the offset frequencies. The resulting bandwidth can approach 100 kHz, as seen below in a magnified view near 18 MHz:

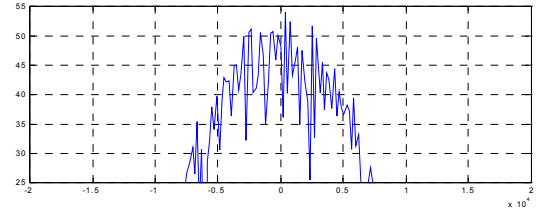


Three CPD Spectral Peaks Near 18 MHz

Derivation of Fine CPD Frequency Structure

Finally, in standard cable plants (STD), the modulators are not locked to a comb generator, and hence do not always produce carriers of exactly the specified frequency. The frequency accuracy specifications for typical modulators are $f_{acc} = \pm 5 \text{ kHz}$ or $\pm 8 \text{ kHz}$ [4]. Hence the actual carrier frequency of any particular modulator will be that specified by the STD frequency plan plus the specified FCC offsets if applicable, and finally plus a very slowly varying random frequency offset selected from a probability distribution with rough limits of either $\pm 5 \text{ kHz}$ or $\pm 8 \text{ kHz}$. Since this is a significant fraction of the medium CPD frequency structure (at increments of 12.5 kHz), one would expect that the result will be a spreading of CPD frequencies about the nominally predicted frequencies, by about half the spacing between CPD medium frequency structure tones.

A separate MATLAB routine was developed to simulate the effect of summing many carriers with random frequencies normally distributed about zero frequency with a standard deviation of 1500 Hz (which gives a 3σ value of 4500 Hz:

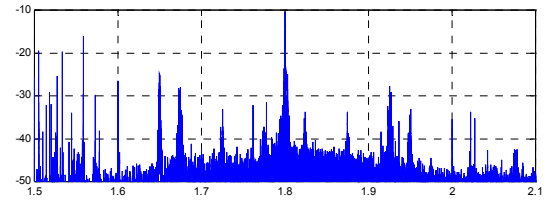


Predicted Fine CPD Spectrum

CPD Measurements

Measured data from a plant using a STD frequency plan and known to have CPD was analyzed and compared to the models described above. The results show complete agreement with the model for coarse, medium, and fine frequency structure, as seen below.

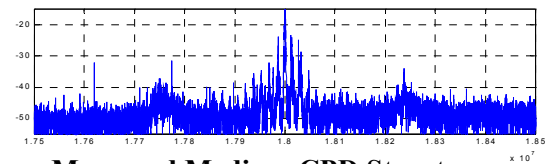
A measurement snapshot of a 6 MHz band centered on 18 MHz is shown below:



Measured CPD Coarse Frequency Structure

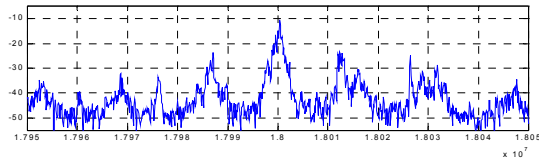
Note the presence of the coarse CPD frequencies at 15.25, 16.25, 16.75, 17.75, 18.0, 18.25, 19.25, 19.5, and 20.75 MHz, in agreement with the model.

The figure below shows the medium CPD frequency structure in the measured data. The two coarse side tones at 17.75 and 18.25 MHz are low enough that it is difficult to discern the medium frequency structure, but at 18.0 MHz, the four strongest medium CPD tones are clearly seen, in complete agreement with the model.

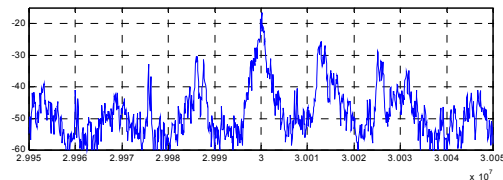


Measured Medium CPD Structure

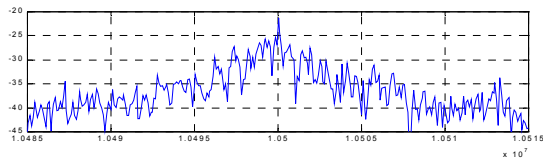
Further detail of the medium and fine CPD structure are shown in the next three figures, where the fine CPD structure can be seen via the non uniform shapes of the medium CPD tones spaced at 12.5 kHz. As discussed above, this is due to frequency inaccuracy in the downstream modulators.



Measured Medium and Fine CPD at 18 MHz

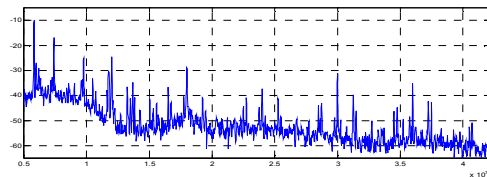


Measured Medium and Fine CPD at 30 MHz



Measured Fine CPD Structure at 10.5 MHz

Finally, a plot was generated of the entire up-stream spectrum from the measured data using the PWELCH function in MATLAB. The resulting spectrum is plotted below:



Entire CPD Spectrum for STD Plant.

Ingress

Ingress of off-air communications has previously been modeled using stationary carriers with Gaussian noise modulation [6]. However actual ingress comes in a variety of forms:

Strong, stationary HF broadcast sources such as Voice of America

Data signals with bursty characteristics

Intermittent push-to-talk voice communications such as ham and citizen's band (CB) radio signals

Slow Scan Amateur TV, allowed anywhere amateur voice is permitted, but usually found at these US frequencies: 7.171 MHz, 14.230 MHz, 14.233 MHz, 21.340 MHz, and 28.680 MHz.

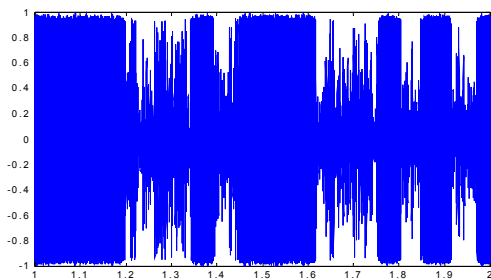
Plus other, less frequent ingress signals such as radar and other military signals

It is relatively straightforward to generate models of all of the above communications using tools such as MATLAB with the signal processing toolbox, since it includes most of the commonly used forms of modulation for such signals. However, the time variation of the signals' power level must be developed. This time variation comes from three main sources: the fluctuations in atmospheric propagation (multipath, ducting, etc.), fluctuations from vehicular movement (in the case of ham and CB radio signals), and fluctuations from the fact that the ingress typically enters the plant in multiple locations. From the evidence that significant reductions in ingress levels occurred after high pass filters were installed throughout the plant, it may be conjectured that ingress typically enters the plant via the subscribers house [7].

From measurements by the author, fluctuations in the signal power of at least 20 dB have been frequently seen with the time scale of fluctuations on the order of tens of milliseconds. Hence, a time varying, random envelope with power variation of up to 20 dB can be impressed on the signals above to generate realistic ingress models for testing new technologies.

For the Morse code communications, captured traces show the on off cycles of such signals to

be on the order of tens of ms ('dots') to hundreds of ms ('dashes'). An example trace is shown on the next page:



One Second of Morse Code Trace

A simple model involves gating a CW signal on and off with a 10 Hz rate to emulate such signals. A more complex model includes specific durations for 'dots and dashes', as well as variations.

Models of voice conversations abound in the telecommunications literature, and can be used for detailed modeling of voice signals. Spaces between words (tens of milliseconds) as well as larger silence intervals (seconds) can be applied to the signal models for single sideband and other common ham and CB voice signals.

The bandwidths of ingress signals range from extremely narrowband on-off keyed Morse code signals, to voice and slow scan TV signals of bandwidth on the order of 20 kHz, to specialized data signals with bandwidths of hundreds of kHz.

The last point in modeling ingress signals is to determine how many ingress signals can occur in band during cable modem signals. This depends highly on whether certain bands are avoided. The DOCSIS specification recommends avoiding the following broadcast bands:

Broadcasting Allocations in 5 - 42 MHz

<u>No.</u>	<u>From</u>	<u>To</u>	<u>Bandwidth</u>
1.	5.95 MHz	6.20 MHz	250 kHz
2.	7.10 MHz	7.30 MHz	200 kHz
3.	9.50 MHz	9.90 MHz	400 kHz
4.	11.65 MHz	12.05 MHz	400 kHz
5.	13.60 MHz	13.80 MHz	200 kHz
6.	15.10 MHz	15.60 MHz	500 kHz

7.	17.55 MHz	17.90 MHz	350 kHz
8.	21.45 MHz	21.85 MHz	400 kHz
9.	25.67 MHz	26.10 MHz	470 kHz

However, this list was developed before advanced TDMA DOCSIS was available. When measurements of upstream ingress are compared with advanced TDMA capabilities, it turns out that many of the above bands have few enough ingressors that advanced TDMA permits cancellation of the ingress. On a typical plant, for example, bands 2, 5, and 7-9 in the previous table can turn out to have relatively few strong ingressors, and thus are candidates. A scan of measured data indicates that with judicious placement of DOCSIS carriers, a typical maximum number of ingress zones to be canceled is about 4-6. The term 'zone' is used, as frequently ingressors occur in groups and must be canceled as a group rather than individually. It is for this reason that the bandwidth to be canceled can often exceed 100 kHz.

Thus, a reasonable model for ingress to use for developing and characterizing advanced TDMA DOCSIS systems is:

4-6 ingressors zones in band

Power levels with up to 20 dB fluctuations over tens to hundreds of ms

Bandwidths in three ranges:

100's of Hz for OOK-CW

20 kHz for voice, data, SSTV, etc.

100 kHz for special signals or for groups of ingress signals to be canceled as a group

Impulse Noise

Random impulse noise has been extensively studied in the past, with the following representing current thinking on the subject [8]:

1 μ s impulse duration (dominant case)

10-50 μ s burst duration (infrequent case)

100 μ s and above (rare)

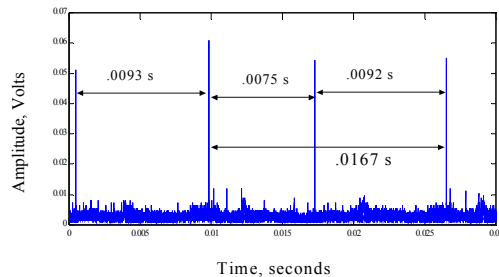
Average interarrival time: 10 ms

As for the spectral characteristics, Kolze [6] gives the following model for random impulse noise:

Each burst event is AM modulated, zero mean, Gaussian noise with a 1 MHz RF bandwidth, carrier frequency of 5-15 MHz (according to measurements) and amplitude ranging from zero to 60 dBmV.

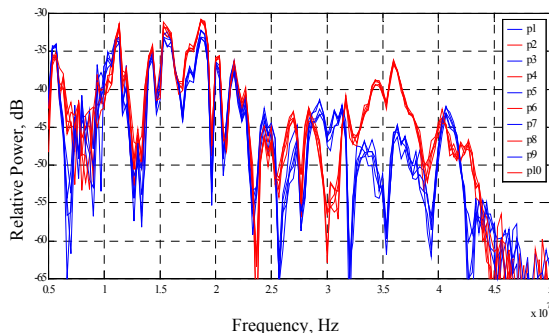
However, periodic impulse noise has been frequently found by the author in measurements on at least one node per headend, and to date no specific models for this phenomenon have been given. These impulses are often quite large in amplitude, and appear to occur with different pulse recurrence frequencies which are usually harmonics of the power line frequency of 60 Hz.

The figure below shows a captured trace of impulses with apparent period of 120 Hz.



Interleaved Periodic Impulses

From the time trace, it appears that the trace is actually two interleaved 60 Hz waveforms. To check, the FFT's of ten successive impulses were captured and plotted below:



FFT's of Successive Periodic Impulses

Clearly, the even impulses are from one element, and the odd impulses from another, perhaps different element. Spectra such as the above are representative of such events captured from multiple nodes and headends.

Hence, in addition to the current models for random impulse noise and periodic noise with period of 60 Hz, 120 Hz, and so on, we should add to the model interleaved periodic trains of 60 Hz recurrence frequency with varying offset intervals between them.

ADAPTATION TO RF IMPAIRMENTS USING DOCSIS ADVANCED TDMA

DOCSIS advanced TDMA provides the operator with several new methods of handling RF impairments. Ingress and CPD may be detected using FFT-based spectrum analysis techniques and then removed using ingress cancellation technologies. Random and periodic impulse noise can be mitigated via greater FEC and interleaving. Both ingress, CPD, and impulse noise can also be mitigated via adapting the modulation order, symbol rate, frequency, and power level. Periodic impulse noise can also be mitigated via detection in the time domain and either avoidance by scheduling around the impairment or mitigation via changing the burst parameters during the expected periodic events. By only changing the burst parameters during the expected impulse events, one can achieve much higher capacities in the upstream than if all upstream packets were forced to use the more robust signaling parameters. Some examples of adaptation are presented below to show how these new features may be used.

Adaptation to CPD

Since CPD has well known and predictable frequencies, it is possible using FFT's of the entire upstream to detect the presence of CPD and move the frequency of upstream signals. The optimum choice would be to move carriers such that the strongest CPD tones (at 6 MHz, 12 MHz, 18 MHz, etc.) lie beyond the edges of DOCSIS upstream carriers. The remaining CPD tones can then be removed by ingress cancellation technology. Note that if CPD is detected, the system should report this event to the operator so that the impairment can eventually be eliminated from

the plant. Until then, the detected CPD frequencies can be maintained in a table, and used for more intelligent frequency hopping if ingress occurs.

Adaptation to Ingress

Ingress cancellation technology makes possible the elimination of ingress from the signal spectrum. Based on the above modeling, a good ingress canceller should be able to cancel 4-6 ingress zones, with occasional cancellation bandwidths of over 100 kHz possible. A key initial decision is whether to notch all ingress frequencies detected by spectrum analysis of the upstream band, or to move the carrier frequency such that strong and/or wideband ingressors are eliminated first from the signal spectrum entirely, leaving fewer and/or more narrowband ingressors to be canceled. The algorithms to make such decisions will depend on the capabilities of each advanced TDMA implementation, and will thus be proprietary in nature.

It is interesting to note that both advanced and conventional DOCSIS TDMA signals can achieve "negative SNR" performance against ingress using the detect and adapt strategy. For example, suppose an advanced TDMA system is using the full 6.4 MHz spectrum for 5.12 Msym/sec signaling and a strong ingressor turns on in band with a power level of +10 dB relative to the cable modem signal. By reducing the symbol rate to 2.56 Msps and moving the carrier frequency to avoid the ingressor, the cable modem system can be considered to operate at 50% capacity using only half of the original 6.4 MHz band where the overall SNR is "-10 dB". This is actually superior to the capacity vs. robustness tradeoff made in a CDMA system where the modem uses the full 6.4 MHz bandwidth, but reduces the number of active codes to provide "negative SNR" operation, albeit at reduced capacity. The difference is that in the TDMA system, the "negative SNR" is almost limitless since the interferer is avoided entirely, whereas the CDMA system spreads the interferer onto the desired signal.

Adaptation to Random Impulse Noise

With the proposed advanced TDMA waveform which incorporates stronger FEC (up to T=16) and RS byte interleaving, mitigation of strong impulse/burst noise is possible, as shown in the table below.

16 QAM @ 1.28 Mbaud, CIR=0 dB		
Interleaver Depth (I)	RS(N,K)=(74,54), T=10	
	Max Impulse/Burst Duration	Max Repetition Rate
I=1	15.6 μ sec	8.6 kHz
I=2	31.3 μ sec	4.3 kHz
I=4	62.5 μ sec	2.3 kHz
64 QAM @ 5.12 Mbaud, CIR=0 dB		
Interleaver Depth (I)	RS(N,K)=(86,54), T=16	
	Max Impulse/Burst Duration	Max Repetition Rate
I=1	4.2 μ sec	45 kHz
I=2	8.3 μ sec	22 kHz
I=4	16.7 μ sec	11 kHz

Adaptation to Periodic Impulses

The periodicity of powerline related impulse noise raises a new possibility for mitigation: scheduling around the impulses. For example, a typical powerline impulse/burst noise waveform might be strong enough to require a low order of modulation such as QPSK, full FEC and interleaving capability, have a repetition frequency of 60 Hz, and a pulse width as high as a few ms. This gives a duty factor (ratio of pulse width to pulse period) of about 1/8. Allowing for error in tracking, an effective duty factor of 1/4 could be used, which means that 25% of upstream minislots would require QPSK with maximal robustness, while the remaining 75% slots could use a much higher order modulation such as 64 QAM. This is possible in advanced TDMA due to the ability to alter modulation order on a burst-by-burst basis. The result would be more than twice the upstream capacity than if QPSK were used on all upstream packets.

CONCLUSION:

APPLICATION TO HFC NETWORKS

The combination of advanced signal processing features, spectrum monitoring, and intelligent algorithms for adaptation in advanced TDMA technology has application to a wide variety of HFC architectures. In large node systems, the robustness provided by advanced TDMA will open up new spectrum that was previously un-useable due to RF impairments. Advanced TDMA permits operation in the presence of impulses, raised noise floor, multiple ingressors, and common path distortion.

On more moderately sized nodes, advanced TDMA will provide a threefold increase in capacity per channel due to doubling the signal bandwidth per channel and 50% more bits/Hz per channel with 64 QAM. In both large and small node systems, the integration of spectrum monitoring, impairment detection and classification, and on-the-fly adaptation of signaling parameters will guarantee that quality of service is maintained for applications such as VoIP.

Finally, in proposed deep fiber architectures such as mini fiber node, the ability to support all of these features in the CMTS means the functionality can be migrated into the mini fiber node itself. In this application, low power consumption, small form factor, and low system complexity are all important to making the mini fiber node architecture cost effective.

References

[1] "Advanced TDMA Proposal for HFC Upstream Transmission," Broadcom Corporation and Texas Instruments, submitted to CableLabs Dec 9, 1999.

[2] *Modern Cable Television Technology: Video, Voice, and Data Communications*, W. Ciciora, J. Farmer, and D. Large, Morgan Kaufmann, San Francisco, CA, pp 608-9.

[3] See for example:
<http://www.fcc.gov/csb/facts/csgen.html>

[4] <http://www.sciatl.com>

[5] Richard Keeler, posting to SCTE List, 3/17/01

[6] *Cable Modems: Current Technologies and Applications*, IEEE Press, Piscataway, NJ, 1999. Part II, The Physical Layer, pp. 135-153, T. Kolze, author.

[7] Email to SCTE-list, Dave Large, 1/04/98

[8] "Analysis of two-way cable system transient impairments," R. Prodan, M. Chelehmal, and T. Williams, NCTA Conference Record 1996.

Contact Information

Daniel Howard, Technical Director
Broadcom Corporation
Email: dhoward@broadcom.com
Phone: 770.232.0018

EFFECTS OF CROSSTALK ON BI-DIRECTIONAL AND HYBRID QAM/DIGITAL DWDM SYSTEMS

Sudhesh Mysore, Ph.D., Director, DWDM Systems Engineering
Oleh Sniezko, Vice President, Engineering

AT&T Broadband, 188 Inverness Drive W.
Denver, CO 80112

Abstract

The increased demand for fiber links in metro market networks has made it necessary for AT&T Broadband to build bi-directional DWDM networks in several locations. Furthermore, the collocation of primary and secondary hub ring routes and the need to address third-party requests for baseband digital capacity have required, at times, the transporting of both analog, or quasi-analog signals such as N-QAM (subcarrier multiplexed signals) and baseband digital signals (SONET, Ethernet or ATM) on different wavelengths within the same fiber. In both cases, the effects of crosstalk generated by DWDM passive components and nonlinear fiber effects had to be analyzed carefully in order to prevent severe degradation of the signal. It became important to specify design rules (relative optical power levels) in DWDM networks and DWDM component performance for proper operation of such networks.

This paper reports on the results of the laboratory testing conducted at AT&T Broadband to determine the acceptable levels of crosstalk, and to specify system guidelines to ensure that crosstalk effects remain within tolerable limits. The paper also summarizes basic specification parameters for DWDM passive components for uni-directional and bi-directional links with analog, quasi-analog and digital baseband signals.

INTRODUCTION

Initially, the metro market systems were supporting video signal transportation from one or a few locations into a number of locations for local distribution. These systems were configured in a ring or a point-to-multipoint topology. Although the optical fiber cable placed for this purpose contained some spare fibers, the fiber count was limited to lower costs, and many of these fibers were used for video signals since DWDM technology had not yet matured.

Later, a number of new services, including local ad signal distribution and insertion, high-speed Internet access, digital telephony, VOD and iTV, competed for the limited fiber capacity in the cable intended primarily for video signal distribution. Additional capacity was also needed to support internal telephony traffic between centralized customer care centers, and from the facilities and the network status monitoring systems to national and regional network operating centers (NOCs). Many of these new services and applications required ring architectures for increased reliability. In many cases, the ring closures were achieved with fibers leased from other operators and the fiber count in these leased runs was limited.

In addition to these internal needs, third party telecommunications service

providers and affiliates who secured agreements with the metro cable operators requested capacity in the same cable routes. Sometimes, the locations they served or their offices were located in proximity to analog fiber links. Moreover, the desire to serve as a CLEC increased the requirements for

digital bandwidth, and often along the analog transport fiber routes. Finally, in some areas, analog and digital transport routes used in different segments of the metro network were co-routed. Some of these are depicted in Figure 1.

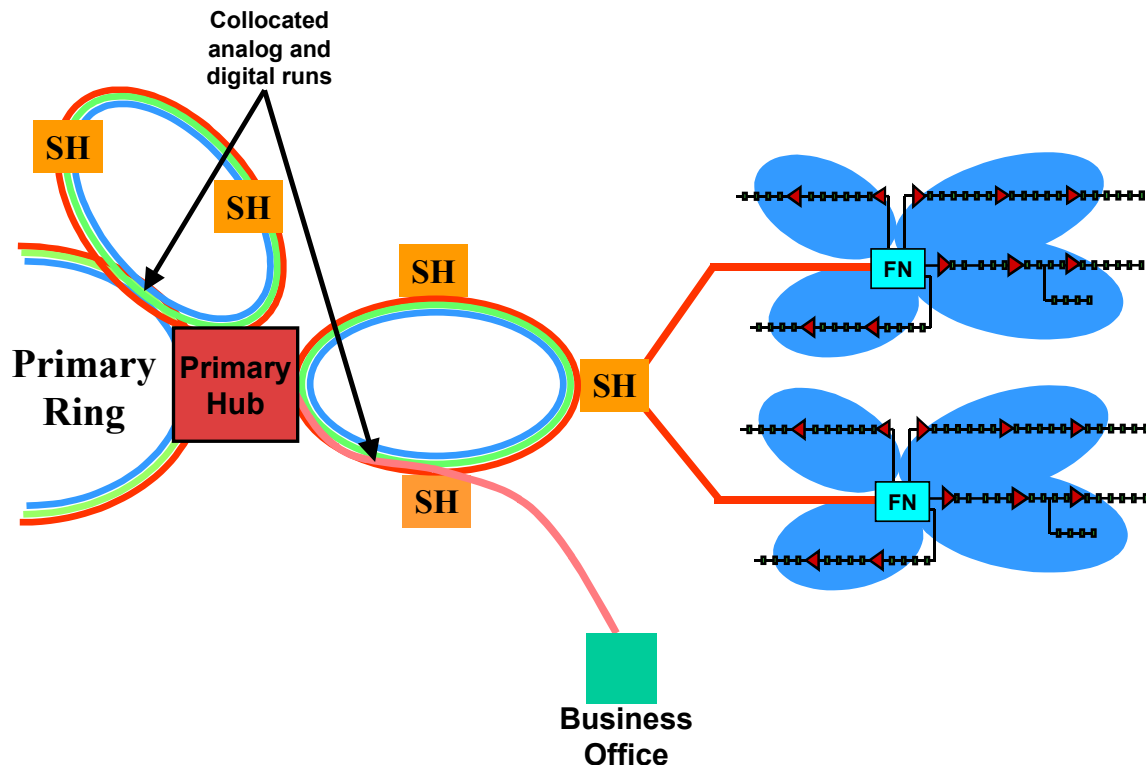


Figure 1: Optical Fiber Network in Metro Markets

These situations made it necessary to use DWDM systems in both primary and secondary hub rings, especially in areas where fiber was scarce or leased. In extreme cases, bi-directional DWDM systems had to be deployed. In a few cases, analog and quasi-analog signals had to be routed over the fiber with baseband digital signals. This situation prompted AT&T Broadband to test, analyze and design DWDM systems that allow for different transport systems on an optical fiber in uni-directional and bi-directional configurations.

The major problems with DWDM links are related to the crosstalk between the

wavelengths used. Although multi-wavelength systems existed before the testing commenced, these systems usually were deployed for uniform transport platforms (for example, all wavelengths supporting SONET transports) and were often integrated with these platforms. In many instances, they were quite expensive due to conservative specifications. The results of the testing and analysis conducted by AT&T Broadband led to a reasonable set of specifications on components from independent sources thus lowering their cost. This approach also resulted in a set of design rules for DWDM systems, dependent on the types of signals in different

wavelengths as well as on the direction of the signal flows. The separation of DWDM systems from the transport systems allowed for moving the legacy systems to multi-wavelength fibers and to recover fibers in some runs.

SOURCES OF CROSSTALK IN DWDM OPTICAL LINKS

Linear Mechanisms

The major linear source of crosstalk in DWDM systems is related to less than perfect isolation of the DWDM demultiplexer. In bi-directional systems, the directivity (i.e., crosstalk from one channel to another) of the DWDM multiplexer is also a critical parameter.

Nonlinear and Hybrid Mechanisms

Nonlinear fiber effects must also be considered if analog signals are transported over a DWDM system. If the optical power coupled in to the fiber exceeds 5 dBm per channel, then Stimulated Raman Scattering (SRS)-induced crosstalk may become a significant¹ contributor to the crosstalk between wavelengths. At these power levels, the worst-case total electrical SRS-induced crosstalk is in the low -50 dB range for the case of a 16-channel DWDM system transporting quasi-analog signals at frequencies above 550 MHz. But, considering that SRS-induced crosstalk is approximately inversely proportional to the RF frequency squared, it could increase to the low -30 dB range (relative to the digital signal levels) at 55.25 MHz (analog channel 2). This is unacceptably high unless the digital signal levels are constrained to be well below that of analog signals.

Other fiber nonlinearities should also be considered. For example, Cross Phase Modulation (XPM) in the fiber results in optical frequency modulation of one signal by the other channels. This occurs because the optical power of one channel modulates the refractive index of the fiber, thereby inducing a phase modulation of all other channels. This nonlinear mechanism combined with a linear mechanism of conversion from phase to intensity modulation results in crosstalk. The XPM is converted to intensity modulation as a result of the non-zero transmission slope of the DWDM filters. This hybrid nonlinear/linear mechanism can be a significant source of crosstalk since a transmission slope of 0.1 dB/GHz represents a 2.3% optical modulation index (OMI) per 1 GHz of frequency modulation.² Typical transmission slopes of demultiplexers are in the 0.02 - 0.11 dB/GHz range within their 1 dB passband.

Another source of crosstalk to consider arises from a combination of XPM and polarization dependent loss (PDL). XPM is known to cause polarization modulation of one signal due to the combined power in the other channels. This polarization modulation is converted to intensity modulation (and hence crosstalk) if the demultiplexer (or receiver) exhibits PDL. The optical crosstalk can increase by 10 dB as the PDL increases from 0.1 dB to 0.5 dB over a frequency range of 50 - 800 MHz.²

EFFECTS OF DIGITAL CROSSTALK ON ANALOG SYSTEMS

A controlled level of either OC-12 or OC-48 signal was optically combined with a 750 MHz analog signal (550 MHz of broadcast and 200 MHz of narrowcast) and

transported over 12 km of singlemode fiber to an optical node. The wavelengths of the analog signal and the digital signal were close together in the 1550 nm window; therefore, there were no correction factors applied for differences in the responsivity of

the receiver. The worst case degradation occurred at lower RF frequencies. A graph of the analog CNR (ch. 2) as a function of OC-12 and OC-48 interfering optical signal level (crosstalk) is presented in Figure 2.

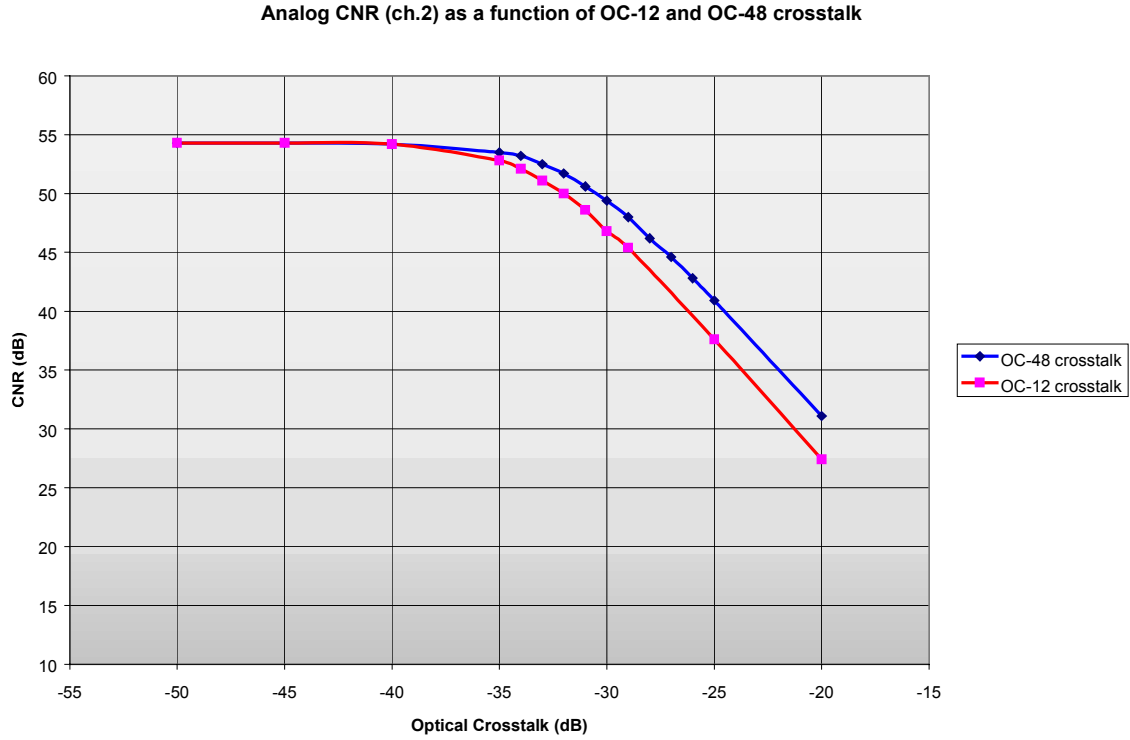


Figure 2: Analog CNR Degradation Resulting from Digital Crosstalk

At high levels of optical crosstalk, the slope of this graph is -2 , indicating that the crosstalk behaves like a simple noise floor. The noise floor is not flat but has a shape given by the spectrum of the digital signal, namely a “sinc” function with the first notch at 622 MHz (for OC-12 signal) or 2.488 GHz (for OC-48 signal). This explains the greater degradation of CNR at lower frequencies and greater degradation caused by OC-12 signal crosstalk. For the same optical power, the spectral density of the OC-12 signal is higher at low

frequencies than the spectral density of the OC-48 signal.

The equivalent CNR of the noise floor resulting from digital crosstalk is given by relationship (1).

$$CNR_{eq} \equiv -10 \bullet \log \left(10^{\frac{CNR_{meas}}{10}} - 10^{-5.435} \right) \quad (1)$$

The equivalent CNR values caused by the OC-12 and OC-48 crosstalk are shown in Figure 3.

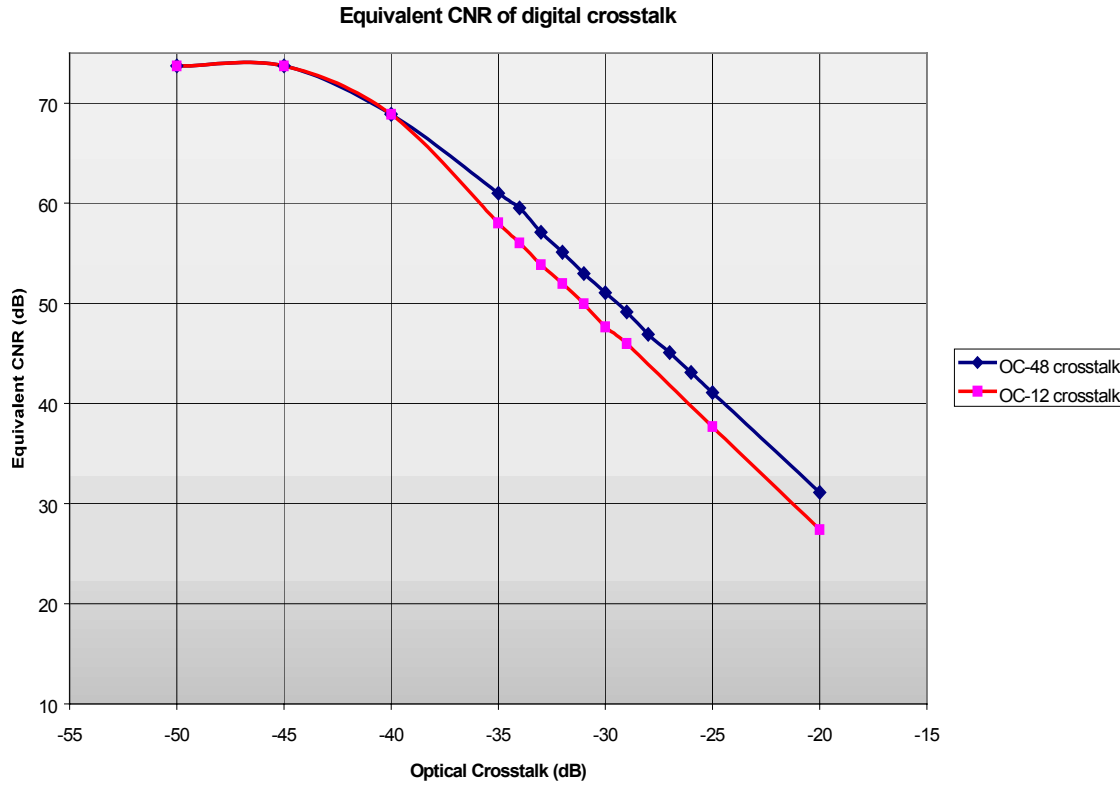


Figure 3: Equivalent CNR Caused by OC-12 and OC-48 Crosstalk Levels

As expected, the equivalent CNR of the OC-12 crosstalk is worse (i.e., lower) than for OC-48 crosstalk. An equivalent CNR figure caused by crosstalk should be higher than 60 dB to ensure that the digital crosstalk has negligible impact on analog performance. Figure 3 indicates that this level of crosstalk can be achieved if the optical crosstalk is lower than -37 dBc. This does not translate into a requirement of 37 dB optical isolation for the DWDM demultiplexer since the digital signals are typically much lower than the analog signals at the demultiplexer input. Typical analog optical power levels at the output of the demultiplexer are 0 to 5 dBm, while digital optical levels are lower than -15 dBm in order to avoid saturating the APD receivers.

At the worst, it can be assumed that the optical levels of the digital channels are at least 10 dB below the analog channels prior to the DWDM demultiplexer. If an EDFA is used at the multiplexer, then optical pads may be required in the digital path (prior to the multiplexer) in order to ensure that the optical delta is not less than 10 dB. Consequently, the required -37 dBc of total digital crosstalk can be obtained using a demultiplexer with a total isolation (i.e., isolation from all wavelengths occupied by digital signals) specification of 27 dB. DWDM demultiplexers with adjacent-channel isolation of 35 dB are available and provide the required total isolation. There are no special requirements for the DWDM multiplexer in this application.

EFFECT OF ANALOG CROSSTALK ON DIGITAL (OC-12 AND OC-48) SYSTEMS

A controlled level of 750 MHz analog signal was optically combined with

an OC-48 signal and transported over 50 km of singlemode fiber to an OC-48 APD receiver. The BER characteristics for the OC-48 system under different test conditions are presented in Figure 4.

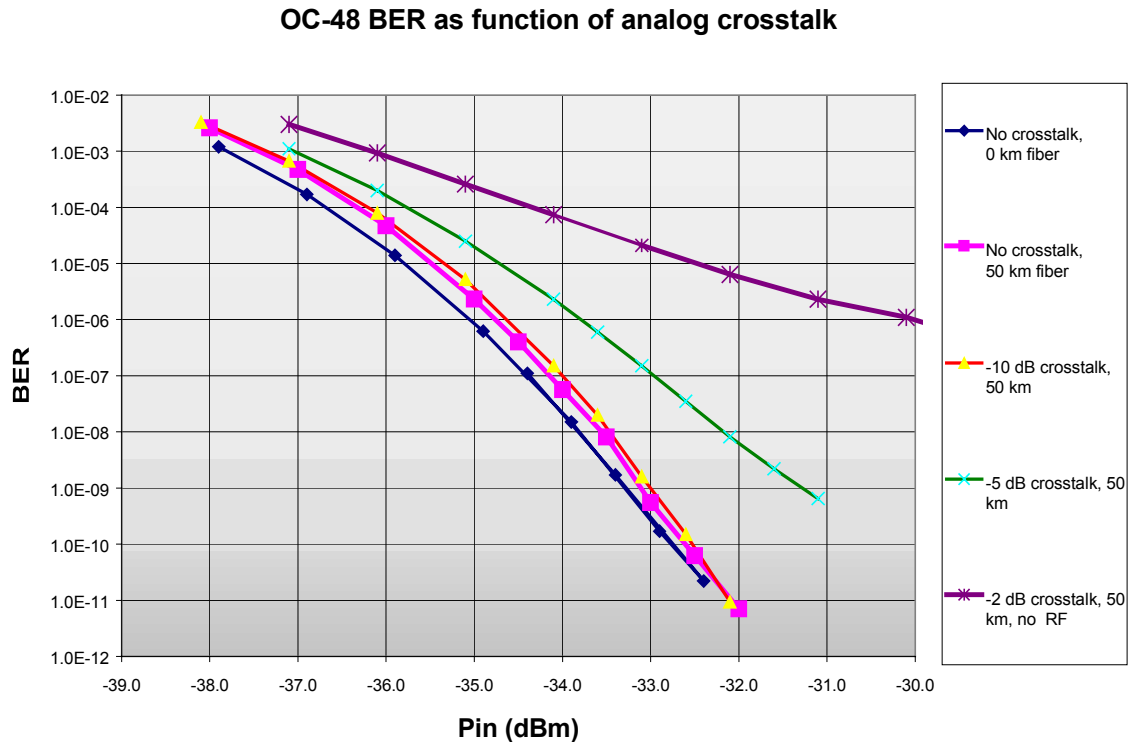


Figure 4: OC-48 BER Characteristics for Different Analog Crosstalk Levels

The leftmost BER characteristic is the baseline characteristic for the case where there is no analog crosstalk and the transmitter and receiver are connected with a 6 m fiber patchcord. To its immediate right is the characteristic for the case where there is no analog crosstalk but there is 50 km of singlemode fiber between the transmitter and receiver. The horizontal displacement of 0.3 dB represents the power penalty resulting from chromatic dispersion. This power penalty would be expected if the spectral width of the laser was 0.13 nm (FWHM).

The next three BER characteristics illustrate the increasing degradation of the digital signal for analog crosstalk levels corresponding to -10 dBc, -5 dBc and -2 dBc, respectively. The degradation becomes significant when the analog crosstalk exceeds -5 dBc. To verify whether this degradation is caused by the crosstalk or by the shot noise generated in the receiver due to the high optical level of the crosstalk signal, BER characteristics for analog crosstalk of -2 dBc were measured with and without RF modulation of the analog source. The results of these tests are plotted in Figure 5.

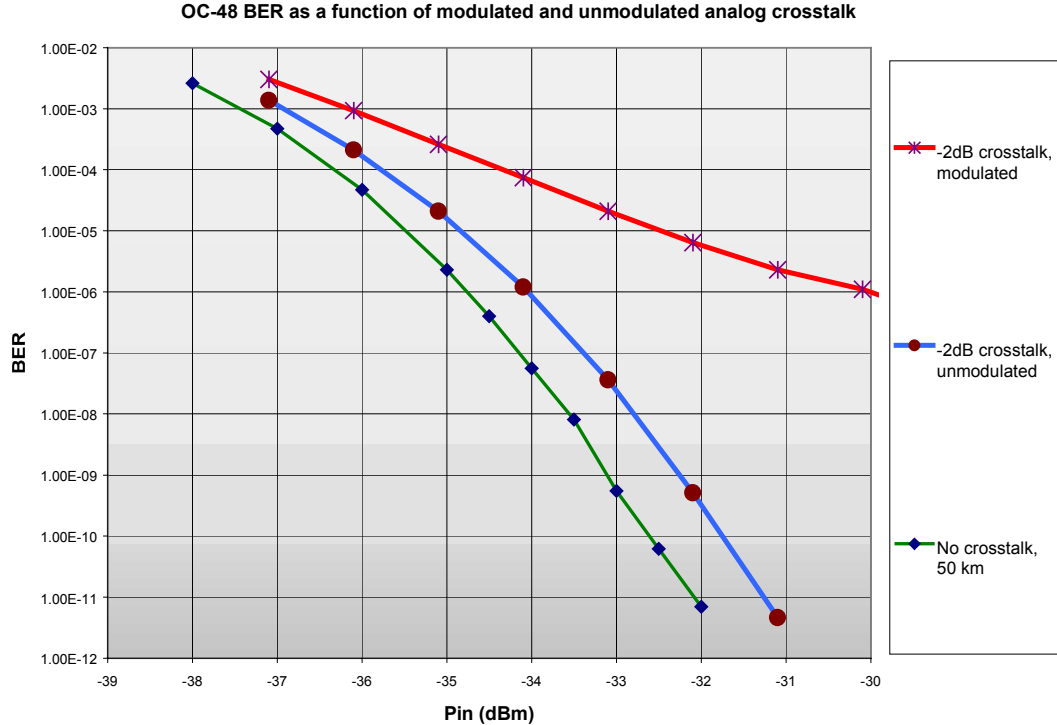


Figure 5: Comparison of BER Degradation Caused by Shot Noise and by RF Interference

The two BER characteristics on the far right correspond to the same average optical crosstalk of -2 dBc; however, in one case the crosstalk signal is modulated while in the other case it is not. The BER characteristic when no crosstalk is present is also included for comparison. Examining the BER characteristic for the case where the optical crosstalk is unmodulated shows that a crosstalk level of -2 dBc is high enough for the shot noise generated by this light level to cause a power penalty of almost 1 dB. However, this is insignificant in comparison to the degradation resulting from the overlap of the digital and analog spectra in the case when the crosstalk light is modulated.

Similar tests with an OC-12 signal showed slightly lower BER degradation resulting from analog crosstalk. This is understandable since there is less overlap

between the analog and digital spectra in this case. Therefore, only the OC-48 results were considered in determining DWDM component requirements.

Based on the test results presented in Figure 4, the acceptable total analog crosstalk should be less than -10 dBc. At this level of crosstalk, the digital power penalty is lower than 0.1 dB. While digital systems can tolerate much higher levels of crosstalk than analog systems (-10 dBc versus -37 dBc respectively), they actually pose a greater design challenge. This is due to the fact that the analog receiver sensitivities are much worse than the digital receiver sensitivities (typically 0 dBm and -30 dBm, respectively). Typically, the analog channel is at 0 dBm and the digital signal is at -30 dBm at the output of a demultiplexer. In order to meet the requirement for the total analog crosstalk to

be lower than -10 dBc, a DWDM demultiplexer with 40 dB of total isolation from all wavelengths with analog load is required. Such a requirement will result in a substantial premium for these passives. Alternatively, the channels adjacent to the digital channels could be left empty. Either alternative would incur a cost penalty.

In systems with higher levels of digital signals (a more typical situation), it is possible to use demultiplexers with an isolation of 30 dB and still obtain a total analog crosstalk lower than -10 dBc by using optical attenuators in the digital paths prior to the multiplexer. The objective is to maintain the optical delta between adjacent digital and analog channels of no higher than 20 dB just prior to the demultiplexer. An optical delta lower than 20 dB and a demultiplexer with total isolation of 30 dB will ensure that the total analog crosstalk is lower than -10 dBc. As previously, there are no special requirements for the DWDM multiplexer.

EFFECTS OF DIGITAL CROSSTALK IN UNI-DIRECTIONAL AND BI- DIRECTIONAL SONET SYSTEMS

The digital systems could tolerate analog crosstalk as high as -10 dBc while suffering a power penalty lower than 0.2 dB. However, they are more sensitive to digital crosstalk. The reason is that, to avoid laser clipping, the optical modulation index (OMI) of subcarrier multiplexed analog systems are much lower than OMI of digital systems. For example, the composite OMI for an analog system is typically 30% (peak) while the OMI for SONET systems is about 90% (corresponding to an extinction ratio of approximately -10 log (5%) = 13 dB). This ratio of three between the OMIs should

theoretically translate into a 5 dB increase in crosstalk sensitivity.

A controlled amount of OC-12 signal was optically combined with an OC-48 signal and transported over 50 km of singlemode fiber to an OC-48 receiver. As in the previous section, BER characteristics were measured for different levels of crosstalk. In contrast to the previous experiments, however, it was not the relative optical crosstalk in dB that was maintained constant for each BER characteristic, but the absolute level of the crosstalk, in dBm. This was done to test a hypothesis that the BER degradation depended on the absolute level of the crosstalk in dBm rather than on the relative optical crosstalk in dBc.

The BER characteristics for crosstalk levels between -60 dBm and -35 dBm are presented in Figure 6. No evidence of the existence of some critical absolute value of crosstalk above which the BER degrades suddenly was discovered. The crosstalk of -45 dBm at the signal level of -30 dBm results in power penalty of less than 0.1 dB. That is, a relative crosstalk of -15 dBc results in a power penalty lower than 0.1 dB for a receiver input level of -30 dBm. At higher input optical levels, the relative crosstalk levels that can be tolerated are higher. This can be ascertained by noting that the horizontal shifts of the BER characteristics are lower than 5 dB if the crosstalk is increased from -45 dBm to -40 dBm or from -40 dBm to -35 dBm. The -15 dBc threshold is used as an acceptable level of digital crosstalk. This is 5 dB lower than the acceptable level of analog crosstalk and agrees with the prediction based on the fact that the OMI for digital system is approximately 5 dB higher than the OMI for analog systems.

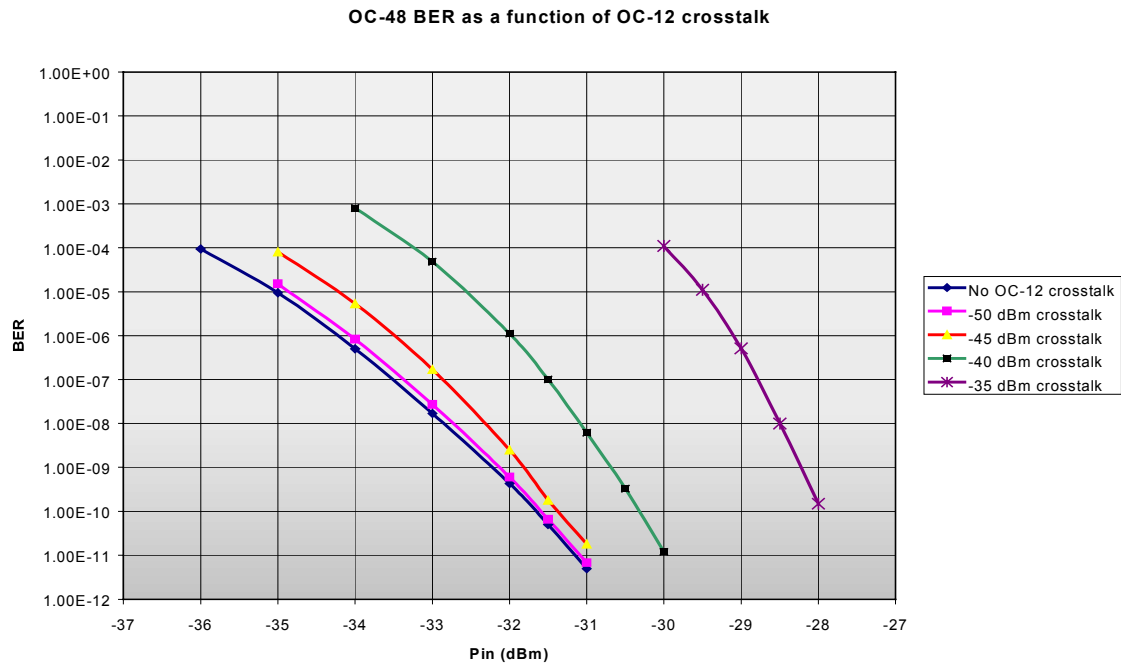


Figure 6: OC-48 BER Characteristics for Different levels of OC-12 Crosstalk

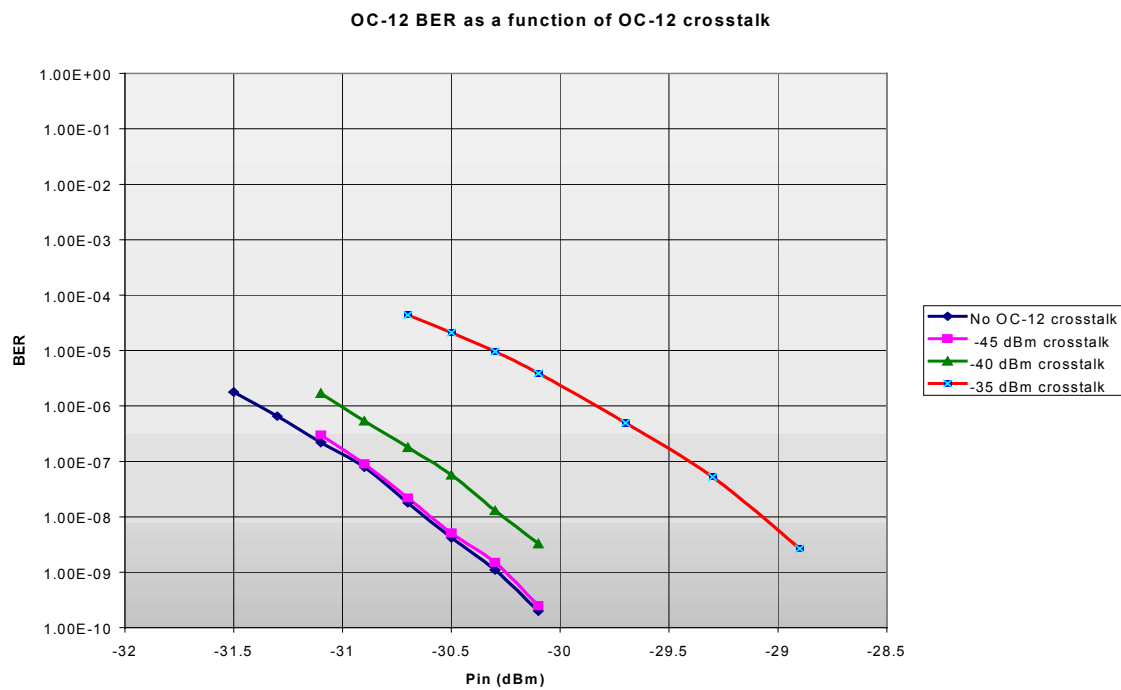


Figure 7: OC-12 BER Characteristics for Different Levels of OC-12 Crosstalk

These tests were repeated to verify the effect of OC-12 crosstalk on another,

independently clocked OC-12 system. The results are presented in Figure 7.

The results indicate that a crosstalk of -15 dBc is still sufficiently low to ensure a power penalty lower than 0.1 dB, even when the crosstalk is from a source with the same SONET rate as the signal. The -15 dBc crosstalk requirement has implications for both the directivity and isolation of DWDM passives in bi-directional SONET systems. A bi-directional system has both transmitters and receivers attached to the DWDM mux/demux at the two ends of the system; thus the mux and demux must have identical specifications. This is in contrast to a uni-directional system where only the demux specifications are of concern in relation to crosstalk penalty and the directivity of the multiplexers is the same as the optical return loss (ORL) requirement of 50 dB. The -15 dBc crosstalk requirement also has implications for the isolation required of DWDM demultiplexers in uni-directional SONET systems.

The allowed amount of coupled light from a strong transmitter back to a port where a receiver is detecting a low light signal can exceed the -15 dBc crosstalk requirement. Assuming the worst case scenario of a $+10$ dBm transmitter and a receiver with a -30 dBm input signal, the directivity of the mux/demux has to be at least 55 dB to keep the crosstalk from the transmitter to the receiver below -15 dBc. This is only slightly higher than the 50 dB directivity of standard (inexpensive) demultiplexers.

The isolation requirements for the DWDM demux depend on the range of received power levels. Since the optical attenuation is approximately equal for all of the wavelengths and the range of transmit powers are fairly well standardized, it can be assumed that the difference between the strongest and weakest received signal is no

higher than 15 dB. Consequently, a total isolation specification of 30 dB is the minimum required to ensure that crosstalk at the receive end is lower than -15 dBc. In order to provide a safety margin, and allow for more stringent requirements for future OC-192 systems, a minimal demultiplexer isolation of 35 dB for both bi-directional and uni-directional applications is recommended.

OC-192 CROSSTALK REQUIREMENTS

For the same optical power, the spectral density of an OC-192 signal is lower (at low frequencies) than the spectral densities of either an OC-48 or OC-12 signal. Consequently, the effect of OC-192 crosstalk on analog systems will be less severe than OC-48 or OC-12 crosstalk. It was shown previously that analog crosstalk affected OC-48 systems more than OC-12 systems. This was explained by the fact that the Nyquist filter in an OC-48 receiver allows more analog crosstalk to pass through than an OC-12 receiver. However, this trend does not continue to OC-192 systems. Since the analog bandwidth is at most 860 MHz and both the OC-48 and OC-192 filters have cutoff frequencies well above this, the effect of analog crosstalk on OC-192 systems should be similar to that on OC-48 systems. Based on previous results indicating that the effects of digital crosstalk on the BER is explained by treating the crosstalk as a noise floor, it is also expected that OC-192 crosstalk of -15 dB or less will have a negligible affect on other digital systems.

It is expected, therefore, that DWDM mux/demux requirements will be unchanged for OC-192 systems. The results of OC-192 experiments to verify these predictions will be presented at the NCTA conference.

SUMMARY OF DWDM MULTIPLEXER/DEMULTIPLEXER REQUIREMENTS

In uni-directional DWDM systems transporting both analog and digital signals, the DWDM mux should have directivity higher than 50 dB and the DWDM demux should have total isolation higher than 30 dB. The optical delta between the analog and digital signals in these systems must be maintained higher than 10 dB (to prevent degradation of the analog signals due to digital crosstalk) but lower than 20 dB (to prevent degradation of the digital signals due to analog crosstalk). This can be achieved by inserting optical attenuators in the path of the digital signals at the mux end, and at the demux end if necessary to bring the received signals within the dynamic range of the receivers.

In uni-directional and bi-directional DWDM systems transporting OC-12 and OC-48 signals, the DWDM demux (and mux in bi-directional systems) should have total isolation higher than 35 dB. Moreover, in bi-directional DWDM systems transporting OC-12 and OC-48 signals, the mux/demux should have a minimum directivity of 55 dB and a minimal total isolation of 35 dB. The maximum transmitted power in the bi-directional systems should be lower than +10 dBm and the maximum spread of received signals in both bi-directional and uni-directional systems should be lower than 15 dB.

In order to avoid problems arising from XPM and PDL, it is also recommended that the maximal transmission slope of a demultiplexer be lower than 0.1 dB/GHz within the passband, and that the PDL be lower than 0.5 dB.

AUTHORS' CONTACT INFORMATION

Sudhesh Mysore, Ph.D., Director, DWDM
Systems Engineering

Email: smysore@broadband.att.com,

Telephone: (303) 858-5204

Oleh Sniezko, Vice President, Engineering

Email: sniezko.oleh@broadband.att.com,

Telephone: (303) 858-5027

¹ M. R. Philips and D. M. Ott, "Crosstalk due to optical fiber nonlinearities in WDM CATV lightwave systems", *J. Lightwave Tech.*, vol. 17, pp1782-92, October 1999.

² M. R. Philips and D. M. Ott, "Crosstalk in cable-TV WDM systems", *Integrated Communications Design Magazine*, pp32-36, September, 2000.

ELIMINATING OPEN ACCESS WOES WITH INTELLIGENT CARRIER-CLASS EDGE ROUTING

Gerry White, Chief Technology Officer
RiverDelta Networks, Inc.

Abstract

Cable operators can embrace Open Access as a wholesale revenue opportunity and as an opportunity to create closer bonds with both residential and corporate subscribers. They can eliminate Open Access woes by deploying intelligent, carrier-class routing at the edge of the cable network to isolate and police individual traffic flows.

Operators can break down traditional barriers to Open Access by implementing carrier-class routing with sophisticated per-flow queuing to support multiple providers of content, applications, and services over shared cable networks. They can use advanced technologies such as MPLS and policy based Routing to deliver end-to-end QoS across the access network and the core networks of multiple revenue-sharing partners.

BREAKING DOWN THE BARRIERS TO OPEN ACCESS

Best-effort data services provide limited revenue growth potential for cable operators. However, by implementing end-to-end Quality of Service (QoS) controls, operators can expand the customer base by offering a wide variety of business and residential services, build increased customer loyalty offering bundled services supporting voice, data, audio, and video traffic, and create multiple revenue streams for the Hybrid Fiber Coax (HFC) network.

However, to fully realize the benefits of Open Access, operators must gain the ability to isolate each traffic flow and police the network infrastructure to ensure that

traffic flows are in compliance with established Service Level Agreements (SLAs).

While Open Access was once viewed as a problem by cable operators, many operators today are realizing the opportunity to accelerate subscriber growth, provide a rich and more complete set of value-added services, and establish profitable revenue agreements with third-party providers.

It is important to carefully define Open Access terminology to understand the technical demands that Open Access imposes on the network. The term “Open Access” means the ability of a cable operator to allow multiple providers to deliver services across the shared cable access network. The term “services” should be interpreted broadly to include content, applications, and other profit-making flows of information.

Operators therefore face the challenge of supporting providers that in the past may have more resembled competitors. But the key to the successful delivery of Open Access is to recognize that the more services are made available to the subscriber, the broader the penetration of cable access networks. Subscribers will select cable as the preferred medium for network services, which in turn increases the total market opportunity for cable operators.

MSOs can continue to deliver their own value-added services, but they will be able to create incremental revenue streams by opening up infrastructure to third-party providers—and gain a percentage of revenue from each new service delivered over the shared network. Open Access does not

trivialize the role of the operator as a mere provider of transport; it creates opportunities for complex and creative business models that enable multiple revenue streams and new opportunities to increase both market share and profits.

Operators need to recognize the diverse business models operators can build to support Open Access. They can deliver IP network services—such as transport, naming, routing, etc.—to enable a basic Internet Service Provider (ISP) service offering. They can also continue to deliver services such as Internet access and Web hosting directly to subscribers.

Operators can create tiered data services to enable Gold, Silver, and Bronze offerings of a given service. This approach allows MSOs to charge premium prices for premium services. Similarly, once they've deployed the technology to support tiered services, they can also allow third-party providers to offer tiered services. This creates opportunities for operators to gain increased wholesale revenues from each service provider partner.

Operators—or their partners—can also deliver enhanced services such as Voice over IP (VoIP), and they can allow Application Service Providers (ASPs) to lease business applications over broadband access networks.

The common denominator of all of these service opportunities is the ability to deploy QoS enabled carrier-class routing at the edge of the broadband access network. Without the ability to isolate and police individual traffic flows, operators lack the control over network resources needed to support multiple providers.

NETWORKING REQUIREMENTS FOR OPEN ACCESS

It is instructive to identify the requirements of the MSO's network within the context of Open Access applications.

In the current Internet access model, a service provider manages IP addresses for subscribers and statically or dynamically allocates unique IP addresses that fall within the address space of that provider. Traffic to the subscriber is then routed to the provider's network based on the IP destination address (which is within the providers address space). Traffic from the subscriber is routed to the desired destination via the provider's network.

Open access will involve assigning IP addresses from the address spaces of multiple providers. These may be delivered from multiple servers or multiple address ranges supplied from a single server. . Traffic to the subscriber can still be routed to the provider's network based on the IP destination address. However routing traffic from the subscriber is more complex, as the path will be dependent on the subscriber's service provider as well as the destination address; e.g. should a packet addressed to a given web site, be routed via ISP1's network or ISP2's network?

The service provider can be determined by the source address of the IP packet so that all the required information is present in the packet. In order to operate in this environment systems must be able to make decisions based on multiple fields in the packet header in real time.

In the Open Access model, services are provided to subscribers from multiple sources. Each provider, therefore, must be able to ensure that their services are working correctly for all subscribers. This is a non trivial problem since each service is based on QoS-enabled IP transport over a shared HFC infrastructure rather than over dedicated

PSTN lines. Effective service management requires MSOs to develop sophisticated QoS and availability parameters and offer third-party providers the abilities to test, quantify, and troubleshoot service delivery of multiple services to all of their subscribers — end-to-end, from the cable modem to the backbone network of each provider.

Quality of Service

Operators require the ability to create and enforce a hierarchy of nested QoS domains within the HFC infrastructure (provider, subscriber, service) which requires sophisticated, high-performance packet filtering and forwarding. Open Access also requires the ability to support end-to-end QoS guarantees across both HFC and third-party networks using industry standards such as Multi-Protocol Label Switching (MPLS) and Diff-Serv.

To provide subscribers and third-party providers with predictable levels of service, it is essential that traffic flows be contained at each level of the QoS hierarchies. Overload or misbehavior within the HFC network by any given provider must be contained within the network resources committed to that service provider — and not be allowed to impact other providers sharing the network. The profitable delivery of Open Access requires advanced isolation functionality to prevent unscrupulous or naive providers from massively overselling their service to the detriment of all other providers on the HFC network.

Similarly, each service provider must be able to isolate each of its subscribers so that none of them can impact other subscribers sharing a common domain. In addition, any overload or misbehavior within a subscriber service should be isolated to that particular service. For example, a CLEC offering Internet access and voice services must be able to prevent a subscriber's web

traffic from impacting that same subscriber's voice calls.

Policing of traffic flows is required to provide the necessary isolation and enable SLA enforcement. Operators need to police traffic flows to make sure that each service provider is compliant with documented SLA parameters. They need the flexibility to ensure that knowledgeable users do not take advantage of the network QoS mechanisms to obtain services for which they have not paid. Traffic that exceeds SLAs should be handled according to SLA policies that determine whether excess flows should be dropped, assigned lower priority levels, or routed at incremental costs.

Carrier-Class Routing

The transition from providing basic Internet access to offering a variety of services from multiple providers moves the MSO from an entertainment provider into a communications carrier. This requires next-generation system that are architected for "carrier-class" reliability, which is usually defined as systems that deliver "99.999%" reliability, which is less than six minutes of unscheduled downtime in a year. Meeting the carrier-class requirements of critical services requires high-levels of redundancy to ensure non-stop operations in the event of a failure of any system component.

Operators must be able to efficiently scale HFC infrastructure to accommodate increased demands for new services and content. This requires next-generation equipment with faster forwarding engines, increased port density, and greater abilities to add network ports so that operators can increase network capacity to support revenue streams from multiple service providers.

As providers aggressively develop partnerships with ISPs and content providers, demand for cable services will escalate. The ability to maximize use of scarce real estate at the distribution hub and regional headend

requires next-generation platforms that provide higher-density RF termination and eliminate the need for external equipment, such as up converters and LAN switches.

Services such as VoIP or streaming multimedia require consistently high-levels of performance, and wire-speed forwarding is required to support a vast array of enhanced services offered by third-party providers. Next-generation, carrier-class edge routing platforms are needed to provide the scalability, density, reliability, and performance needed to support Open Access. Operators need to be able to ensure that carrier-class platforms deliver the guaranteed SLA requirements that they have committed to both provider partners and to subscribers.

Service Provider Selection

A subscriber should be able to select from multiple providers based on the competitive nature of their offerings, such as Internet access from a selection of ISPs, video service from the MSO, and voice service from Competitive Local Exchange Carriers (CLECs) or InterExchange Carriers (IXCs). Both residential and corporate customers should be able to select services either on a subscription or pay-per-use basis. This requires flexible, open systems provisioning and management combined with sophisticated, high-performance routing.

Operators need to support advanced SLA parameters such as maximum bandwidth allocation, minimum bandwidth guarantees, bounded delays, and bounded jitter. They will need the ability to define QoS parameters both statically (e.g., Gold/Silver/Bronze services) and dynamically (e.g., for services such as voice call set-up). At a minimum, operators need the QoS capabilities of DOCSIS 1.1-based equipment, but they also need features beyond these standards to enable enhanced services over both HFC and service provider networks.

Metering/Billing/Reconciliation

Allowing multiple service providers to operate over a shared access network requires robust features for reconciliation and billing. Detailed accounting information needs to be maintained on a per-flow basis to ensure that SLAs are enforced, and the sophistication and complexity of accounting can vary dramatically.

In the simplest case, a provider could define an SLA and the MSO could implement a policing mechanism to ensure that it is not exceeded. However, in most applications both the provider and operator will want to meter the SLA to ensure conformance. If subscribers have access to pay-per-use services such as long-distance phone calls or videoconferences, then the MSO needs to offer metering services that can support dynamic billing. Billing models based on both time-of-use and traffic volume is required with an event-driven mechanism used to initiate and terminate metering at wire speed.

OPTIONS TO PROVIDE OPEN ACCESS

In theory, an MSO could create and maintain multiple RF channels to carry traffic for each provider. Lack of sufficient RF frequencies and the requirement to duplicate CMTS systems per provider render such a solution impractical.

Fortunately more viable alternatives are available. These can be classified into two general categories. Tunnel based solutions in which subscribers are tunneled back to a centralized subscriber management platform responsible for implementing traffic policies and routing subscribers to the appropriate provider networks.

Policy-based routing solutions in which the edge router/CMTS system is responsible for implementing traffic policies and for routing subscriber traffic to the appropriate provider network.

TUNNELING :A CIRCUIT BASED APPROACH

Generally, tunneling is used for dial-up Internet and DSL access. Subscribers connect to a network access server using a modem connected to the public switched telephone network (PSTN) or a DSL circuit. In these networks, a subscriber management system located inside the network manages the traffic flows.

Traffic flows reach the subscriber management system via a tunnel mechanism such as a Point-to-Point Protocol over Ethernet (PPPoE) or Layer Two Tunneling Protocol (L2TP) tunnel built on top of the generic network infrastructure. Once the flow reaches the subscriber management system, the system terminates the tunnel, examines the data received, implements QoS and policing and directs the traffic flow to the required application server.

This mechanism requires client software on the host system to initiate the subscriber end of the tunnel, which can present an ongoing support problem. The most serious drawback to tunneling is that it hides the content of the flow. Because the CMTS cannot recognize what the tunnel carries, the HFC access network cannot use the QoS built into DOCSIS 1.1. Application-based QoS is not available to traffic within the tunnel. Without the ability to give voice or video traffic higher priority, operators will have difficulty meeting the performance guarantees promised for these services.

The “bandwidth tax” associated with tunneling is also significant. Tunneling requires additional headers on top of the DOCSIS protocol. This approach wastes bandwidth in the access segment, where network capacity is most strained.

Because tunneling requires that the subscriber management system be located inside the network, operators must place applications servers even deeper in the

network. Thus they negate the benefits gained from moving content for high-bandwidth services closer to the user.

Finally, tunneling deprives the cable operator of one of its most powerful weapons against its DSL competitors--the “always-on” connection. Before a user can have access to even the most basic e-mail services, the tunnel must be established. It is therefore difficult to deliver push services like newscasts. Likewise, multicast services are problematic, because MSOs must convert a multicast into multiple unicast messages at the subscriber management system, which further hogs scarce bandwidth.

POLICY-BASED ROUTING

Policy-based routing differs in that the router manages traffic flows at the edge of the network. The router looks at multiple fields within packets to determine the appropriate routing and QoS.

Each user is provided with an IP address in an address scope associated with his/her selected service provider. Packet routing is partially determined by looking at the source IP address, understanding to which service provider partner the IP address belongs, and then routing the traffic to that partner for their handling. Such examination allows the router to implement more sophisticated QoS policies than are possible by simply looking at the data’s destination address.

Two variants of policy based routing can be considered for cable networks, a centralized model with the policy router located at the regional head end (or other convenient central location) and a distributed model in which the policy routing function is moved to the edges of the network (e.g to a distribution hub).

Both distributed and centralized architectures require a DHCP address management system and a policy based router. The distributed solution places the

policy router with the CMTS and uses an MPLS based metro or wide area network to connect the policy routers to the service provider networks.

DHCP Server

A DHCP server provides each user with an IP address in an address scope associated with his/her selected service provider. This address is used to identify the service provider to which the customer has subscribed. The DHCP server is well understood technology and need not be described further.

MPLS Virtual Networks

MultiProtocol Label Switching (MPLS) is a standard under development by the IETF for efficiently switching IP traffic over IP or ATM core networks. MPLS adds a label to IP packets which instructs network routers and switches where and how to forward the packets.

Today's conventional routers analyze IP packets at each hop in the network, which is a time-consuming process. With MPLS, an intelligent edge router (Label Edge Router or LER in MPLS terminology) looks at the header of the first packet in the traffic flow. Based on the header's contents, the router applies a label to that packet and all subsequent packets in the flow. This label determines where to send the entire data stream and the QoS policies to apply. For example, the label may indicate that the flow contains a Voice over IP (VoIP) call destined for a particular voice service provider. The label would dictate that every packet be placed in a low-delay path with guaranteed delivery to the provider POP. This path is known as a Label Switched Path or LSP. The routers or switches in the core of the MPLS network (Label Switch Routers or LSRs) do not examine the IP headers of the packets but

instead switch the packets based on the appended MPLS labels.

Labeling packets at the network's edge eliminates the processing traditionally performed in the core. The core network can therefore focus on switching traffic to its destination as quickly and efficiently as possible. Bottlenecks in the core are reduced or even eliminated.

MPLS enables cable operators to offer a variety of services over shared network infrastructure with differing QoS requirements without overloading the core network with unnecessary processing.

For a distributed Open Access architecture a virtual network is created for each service provider between the service provider points of presence (POPs) and the edge router / CMTS at the distribution hubs. This virtual network is based on MPLS technology and consists of a mesh of MPLS label switched paths (LSP's) which are set up between the CMTS and the POP. Each LSP provides a traffic engineered path over the shared metro or wide area network transport.

CMTS / ER

The CMTS/ER located at the distribution hub looks at multiple fields within packets to determine the appropriate routing and QoS. Packet routing is partially determined by looking at the source IP address, understanding to which service provider partner the IP address belongs, and then routing the traffic to a specific LSP in the virtual network of the partner for their handling. Such examination allows the router to implement more sophisticated QoS policies than are possible by simply looking at the data's destination. The CMTS/ER is functioning as an MPLS label edge router in this case.

The policy information required by the edge router to enable these complex forwarding decisions can be disseminated

using policy extensions to existing routing protocols such as OSPF and BGP4. BGP4 is the preferred routing protocol used for connection between autonomous systems and as such is a favoured approach for the policy router.

For policy-based routing to work in a cable environment, operators need to deploy a high-powered QoS capable routing engine used in combination with a DOCSIS 1.1-capable CMTS in the distribution hub. The CMTS/ER should enable QoS to be maintained from the cable modem to the service provider POP. Thus it must implement QoS on the HFC network by mapping IP flows to DOCSIS 1.1 service flows; it must maintain the QoS through the CMTS / router and then map these flows to MPLS label switching paths with the desired QoS characteristics.

Upstream packets from the HFC network are scheduled by the CMTS/edge router according to DOCSIS 1.1. First, the packets are classified by the cable modem, which requests transmission on the appropriate DOCSIS flow. At the CMTS, the packets are re-classified based on filters and then QoS policy can be applied. Each flow can be assigned its own queue to ensure QoS is maintained through the CMTS /ER, and then packets are forwarded to the LSP required to reach the network of service provider partners based on fields in the packet header such as the source IP address.

Downstream packets are received from the (traffic engineered) LSP, mapped into downstream flows based on the IP header fields and scheduled for transmission onto the DOCSIS downstream channel. By providing sophisticated queuing and scheduling mechanisms QoS may be maintained in this direction also.

Advantages of Policy-Based Routing

Unlike tunneling, policy-based routing is designed for a broadband IP infrastructure. Because the router can look at the individual application flows, it can extend the capabilities of DOCSIS 1.1. It can assign QoS and routing policies based on parameters such as service provider, subscriber, and application. Per-flow queuing enables the router to isolate the traffic of different services and different providers at the edge of the DOCSIS network.

Because policy-based routing conforms to IP standards, current features like transparent IP multicast can be supported. Policy-based routing can also take advantage of future developments in IP standards such as the rollout of MPLS.

DISTRIBUTED vs CENTRALIZED POLICY ROUTING

As operators evaluate the networking requirements for Open Access, they should consider whether the policy routing intelligence should reside at the edge of the network or at a central location.

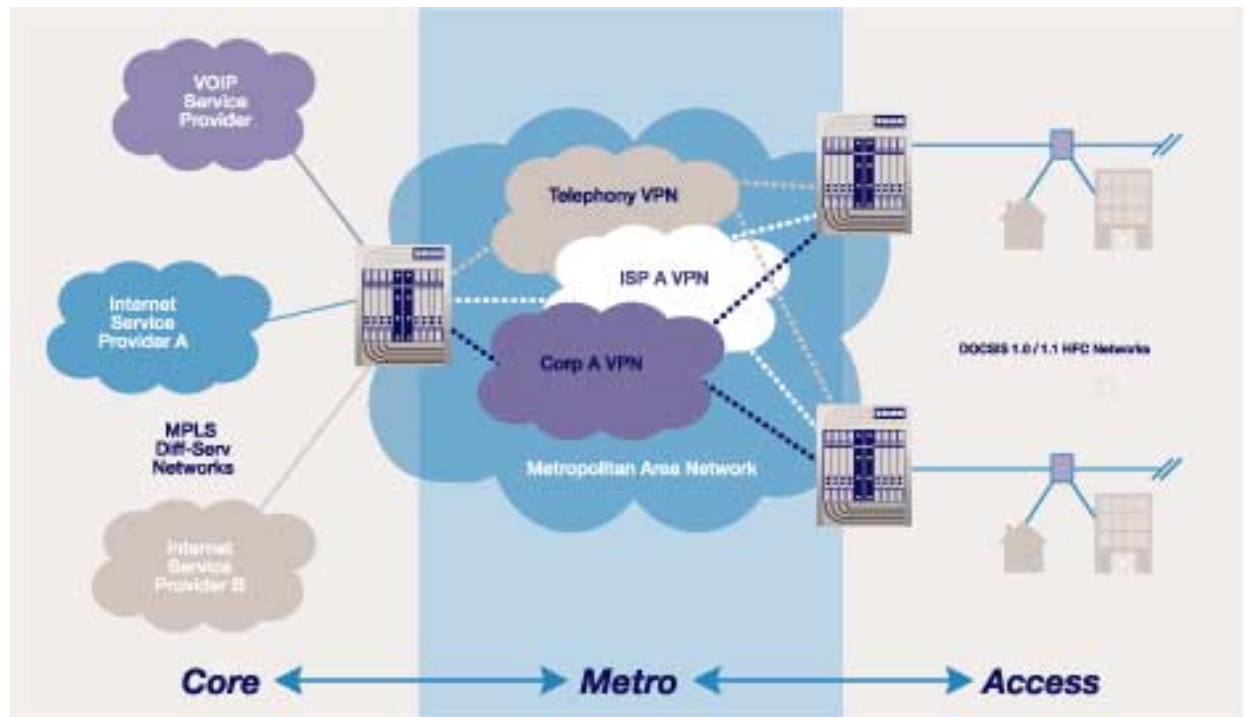
A centralized solution leads to inefficient bandwidth use. All traffic must be routed to the policy router for treatment—no matter its destination. As subscriber penetration of new services increases, this approach can cause bottlenecks. With intelligent edge routers, operators can contain local traffic on the network and establish the optimum routing path for internetwork traffic.

Pushing intelligence to the edge also allows operators to move application and content servers closer to subscribers. This is a plus for customers because they'll see increased performance for which both operators and partners can realize premium pricing.

The distributed model has much better scaling properties as the work intensive policy decision making is distributed. It also simplifies provisioning multiple connections to provider networks from the MSO regional

network to support redundancy and load sharing..

A distributed solution is shown in the figure below.



Distributed Policy Based Routers with MPLS based Virtual Networks

REQUIREMENTS FOR A POLICY ROUTER AT THE NETWORK EDGE

Operators need to be able to deploy intelligent policy based routers at the edge of the network to isolate and police traffic flows. These next-generation edge routers need advanced intelligence so they can classify, manage, and police the traffic from the cable modem to the core networks of service provider partners. Both third-party providers and subscribers will demand SLAs based on guaranteed QoS levels, and edge routers must be capable of implementing these QoS guarantees end-to-end.

Carrier-class routing protocol implementations are essential because the edge router should be capable of peering with other ISPs to make Open Access a reality. A full suite of unicast and multicast routing protocols is required to allow interoperation with peer routers, (e.g. RIP v1; RIP v2, OSPF v2, BGP4, IS-IS, DVMRP, PIM-SM). Providing a carrier grade implementation of the routing protocols requires more than a basic implementation of the minimal functions required for conformance. As well as offering a full feature set it must provide a robust, highly available solution which is resilient to network errors and attacks The

routing software implementation must scale in terms of numbers of routes, interfaces, and peering relationships to support expansion as new services, subscribers, and providers are added. It must also provide support for operation staff to detect and resolve network routing issues.

The edge router must apply policy functions on a per-flow basis and must be able to provide a guaranteed minimum rate per flow to enforce SLA commitments.

The router must isolate the traffic of individual providers, subscribers, and applications. When IP traffic from a provider exceeds its SLA, the edge router implements predetermined policing policies to ensure that each flow receives at least its minimum guaranteed bandwidth. When congestion occurs, the router should drop packets from misbehaving flows instead of dropping traffic that is operating within its SLA.

Operators will need sophisticated accounting and metering systems in a multiprovider environment because wholesale providers will of course require proof that they are receiving committed performance levels. Again, an intelligent edge router can provide these statistics. Because operators will be peering with other ISPs and service providers, it is critical that edge routers conform to the highest standards of interoperability. Partners will require the ability to control their portion of the network; so flexible standards based network management is essential.

Implementations require a high-powered QoS routing engine and a DOCSIS 1.1-compliant CMTS. QoS on the HFC network is provided by mapping IP flows to DOCSIS 1.1 service flows based on contracted service levels, and the QoS on the metropolitan network is provided by mapping IP flows to traffic engineered label switched paths in the MPLS networks. Thus cable operators can guarantee performance to their

wholesale partners and generate additional revenue streams.

Edge routers are the transition point from the HFC access network to the regional backbone. They identify and classify traffic flows, apply QoS, implement admission control and efficiently forward traffic to its destination—which can be the core network of one of multiple providers.

By applying this intelligence at the edge of the cable network, MSOs can provide end-to-end QoS across the HFC network and across the backbones of multiple revenue-sharing partners.

AVOIDING OPEN ACCESS WOES

The use of hierarchical per-flow queuing and carrier-class edge routers allows cable operators to benefit from Open Access and ensure maximum control over network resources. By selecting next-generation edge router/CMTS platforms, operators can welcome Open Access as an opportunity for new revenue streams.

They will be able to bind subscribers to the cable network, improve bandwidth utilization, and increase profits. Operators will be able to classify and treat individual traffic flows and deliver QoS guarantees across access, metropolitan, and core networks, and they will be able to allow multiple revenue-sharing partners to offer diverse portfolios of services that will create tighter loyalty to cable networks for both residential and corporate subscribers.

Policy-based routing has not been widely deployed up to this time since legacy routers lack the performance, scalability, and per-flow processing necessary to implement it effectively. They can't perform source based, content-aware routing because the per-flow packet classification and QoS required is beyond their processing capacity. But with today's high-powered silicon and advances in QoS theory, next-generation edge routers can

examine individual traffic flows and provide the forwarding and QoS functions required at wire speed.

ABOUT THE AUTHOR

Gerry White is Vice President and Chief Technical Officer for RiverDelta Networks, which designs, develops, and markets intelligent edge routers for broadband access to optical networks. He is the co-author of several patents and articles on data communications technology and he can be reached at gwhite@riverdelta.com. RiverDelta Networks can be found on the Web at www.riverdelta.com.

ENCODING AND DELIVERY OF MPEG-2 COMPRESSED CONTENT TO CABLE HEADENDS

John Vartanian
iN Demand

Abstract

Many technological advances are making video-on-demand (VOD) economically feasible for cable operators. Fiber rich architectures, the deployment of digital settops and dropping file server costs have resulted in virtually all of the major MSOs planning to either test or deploy VOD on their systems in the near future. Although capital costs have been dropping rapidly, operational expenses are still significant. The cost of encoding, distributing and loading content onto VOD file servers is a major expense if borne by a single operator.

One way to reduce these operational expenses is to spread costs across many systems offering video-on-demand. To this end, iN Demand has embarked on an effort to encode content to a single, open specification that works with a variety of file servers and settops, and then distribute that programming through a nationwide satellite network.

Encoding

Encoding content is a significant expense, especially when the number of potential VOD subscribers is small or the programming is not expected to generate many VOD purchases. It is a labor-intensive process that requires the expertise of a compressionist. By spreading the fixed encoding cost across a wide base, this expense becomes less burdensome to individual operators.

An encoding specification is necessary to ensure that content is compatible with the various file servers and settops that will be used to deliver VOD. Many of the details of this specification deal with packet identifier (PID) definitions, packet lengths, program clock reference (PCR) information and group-of-picture (GOP) structure. All of these are important and necessary to ensure the interoperability of content. However, I will focus on three areas that are more discretionary: video bit rate, resolution and audio.

Many encoding decisions affect the economics of offering VOD. Bit rates result in tradeoffs that determine the number of

simultaneous users on a VOD network. Video resolution ensures high quality by properly allocating bits between detail and motion needs. Provisions for high quality audio mean that VOD delivered by cable operators will not be at a competitive disadvantage to DVDs.

The first step in encoding a movie is to acquire a dub of the movie master. In order to minimize video artifacts, dubs for digital video compression are usually provided in a Digital Betacam or D1 format. Both of these digital component formats allow the compressionist to start with the best possible source material.

Bit Rate

iN Demand encodes content that is MPEG-2 compliant, but with additional parameters defined to meet the needs of VOD on cable systems. One of the most important parameters in the specification is the video bit rate assigned to each VOD stream. Bit rates that are too high sacrifice VOD availability without a visible improvement in quality. Bit rates that are too low result in quality degradation that can be perceived by VOD users.

The video bit rate in the specification was determined by holding subjective tests with a dozen viewers, including representatives from the technical, marketing and business development areas. A five-minute section of *The Matrix*, which was a challenging movie to encode, was compressed with varying parameters. The clip was encoded at video bit rates that changed in 100 kbps increments from 2.6 Mbps to 3.4 Mbps. Compression was done for each bit rate at full, $\frac{3}{4}$ and $\frac{1}{2}$ resolution. In addition, compression for each bit rate and resolution was done twice: once at a constant bit rate and then at a variable bit rate, but with the CBR set as the maximum allowable bit rate.

In all, the bit rate, resolution and CBR/VBR permutations resulted in 54 versions of the test clip.

The consensus of the evaluators was that a video bit rate around 3.1 Mbps was needed to produce good video quality with material that is challenging to encode. Bit rates below that resulted in perceived degradation of the video quality. Bit rates above that did not increase the video quality.

We found that CBR produced better video quality than VBR within the constraints we set on the VBR clips. While the advantages that true VBR bring are desirable (as seen with DVD encoding), it is not economically feasible at this time to process VBR streams in a headend without reducing the number of streams carried in a cable channel. It is anticipated that once rate remultiplexing and other video processing technologies are further developed, VBR will lead to average bit rate reductions without degradation of video quality or a reduction in VOD system capacity.

Resolution

One result of the evaluation that was somewhat surprising occurred in the area of video resolution. With full resolution, all the picture elements, or pixels, are evaluated while encoding a movie. In this mode, the encoder uses a frame of video that is typically 720 pixels in a horizontal line and 480 active vertical lines. In the $\frac{3}{4}$ resolution mode, the encoder reduces the number of horizontal pixels to 544 or 528, but still keeps the 480 vertical lines. With $\frac{1}{2}$ resolution, 352 pixels across and 480 lines down are used. Although it may seem that using full resolution would provide the best quality, this was not the case. Since a consumer television cannot display all the detail encoded at full resolution, it is inefficient to encode and transmit this extra

data. Instead, either $\frac{3}{4}$ or $\frac{1}{2}$ resolution provides higher quality. In these modes, bits that would be used to encode detail in the full resolution mode are instead used to encode motion. The result is no loss of detail on a consumer television and a reduction in motion artifacts that would occur at full resolution. In general, we found that $\frac{3}{4}$ resolution usually provided the highest quality, but $\frac{1}{2}$ resolution was better in some scenes that contained a lot of action.

Audio

In addition to the video requirements, a portion of the bit stream must be reserved for several other components. Dolby Digital audio, which provides state-of-the-art surround sound for home theaters, requires .384 Mbps. This type of audio is the standard for DVDs and some DBS channels and must be included in order for VOD to be competitive with these other technologies.

We also looked at secondary language support. If alternate language tracks are multiplexed with the video and English audio, additional bandwidth needs to be reserved. Instead, the specification calls for additional audio services to be assigned their own audio PIDs. The desired audio is multiplexed with the video at the file server and only that one language is delivered to the consumer settop. This method of multiplexing allows for a virtually unlimited number of language choices, with no effect on bandwidth requirements.

QAM Loading

Once we determined the minimum video bit rate and the audio requirements, we looked at QAM loading. Program specific information (PSI) and overhead needs were added to the 3.1 Mbps video and the .384 Mbps audio. This resulted in a minimum

need of approximately 3.67 Mbps per stream. At this rate, seven streams can fit into the 26.97 Mbps payload of a 64 QAM channel. The maximum number in the 38.8 Mbps payload of a 256 QAM channel is 10 streams. Once the QAM channel loading was determined, the video bit rate was increased from 3.1 Mbps to 3.18 Mbps. This allowed for a higher video bit rate, without a decrease in the number of streams in a 64 or 256 QAM modulator. After some other small requirements are added, a total multiplexed rate of 3.75 Mbps per stream became the bit rate specification.

Meta Data

Meta data is any information that is associated with a movie. Common meta data elements for a movie include its length, a summary, the actors, a rating and the movie's availability dates. Meta data is necessary in the operation of a VOD service. It allows a customer to do such things as set parental lockouts or sort a movie by genre.

As with content encoding, it is efficient to add meta data to a movie once, with the meta data being interoperable with different VOD applications developed by various file server and settop vendors. This removes the need for insertion of meta data at each VOD system. In addition to reducing labor in the headend, centralizing meta data functions reduces the opportunity for errors that could occur if meta data had to be entered at each VOD system. iN Demand is using a meta data standard that defines each meta data field, its type and its length.

Efforts are underway to expand the meta data standard to incorporate all meta data needs by the major VOD file server and settop vendors. By creating a superset of meta data and sending the same version to everyone, systems get the advantages of

centralized meta data creation. Except for movie promotion, which includes an MPEG-2 video trailer and a JPEG or GIF box cover or poster art image, meta data is text information that does not require significant additional bandwidth. There is little disadvantage to sending a superset of meta data to all systems. Each VOD application at a cable system can pick the meta data fields it needs and ignore the others.

Distribution

Once a movie has been encoded and meta data has been added, the next challenge is to distribute the asset package to VOD systems. The easiest way to distribute compressed movies, at least in the early stages, is by sending digital linear tapes (DLTs) via an overnight shipping service. However, this method of distribution becomes inefficient when the number of deployed VOD sites grows beyond a few. Satellite delivery is a natural means for sending content to cable headends. Once the satellite distribution network has been built, it is easily scalable. Additional VOD systems can be added at a relatively low incremental cost.

Satellite Link

iN Demand distributes pay-per-view movies, events and sports on eight 36 MHz C-band transponders. To maintain compatibility with cable systems that are using 64 QAM to deliver digital signals, 24 MHz of each transponder is used for MCPC transmission. The other 12 MHz in the transponder creates an opportunity for VOD file transmission.

The first step in our satellite analysis was to perform a link budget to determine the power level of each carrier and their center frequencies. When multiple carriers occupy

a transponder, the power of each carrier must be reduced to minimize intermodulation products. Once the power levels and carrier frequencies were determined, the MCPC carrier was moved from the center of the transponder to its edge, allowing the maximum bandwidth for the new VOD carrier. An 8 Mbps carrier was then added to the unoccupied transponder bandwidth.

IP Multicast System

On the following page is a diagram of the IP multicast system that iN Demand is using deliver VOD files to cable systems. The pitcher is at the heart of the IP multicast system. It is the component, along with external hard drives, that stores the assets and sends them to systems using an IP multicast transmission. It also performs a triple DES encryption, ensuring the security of the VOD file.

The pitcher is controlled by the manager, which schedules the delivery of assets to headends. It analyzes available satellite bandwidth, asset size and required delivery dates to determine when assets need to be transmitted by the pitcher. The manager is operated through a user interface on a remote terminal, with information sent via an Internet connection from the terminal to the manager. The manager also provides monitoring and management functions, both of which can also be accessed at the remote terminal.

An asset creation tool (ACT) adds meta data in an XML format to the movie. It then transfers the content and associated meta data to either the pitcher's internal or external storage, where it resides until the pitcher receives transfer instructions from the manager. As with the manager, the ACT is also accessed through a user interface on a

remote terminal, with information sent to the ACT over the Internet.

When the manager instructs the pitcher to transmit a file to a group of systems, the compressed content is sent to an IP encapsulator. The IP encapsulator creates an MPEG transport stream from the IP data and translates the data into a form suitable for satellite delivery. It also applies a $\frac{3}{4}$ rate forward error correction to compensate for transmission errors in the satellite link. The MPEG transport stream is then sent in an ASI format to a quadrature phase shift key modulator and uplinked to the satellite.

The VOD carrier is downlinked at cable headends and sent via L-band from the satellite antenna to a catcher, which is a Pentium III device operating on Linux. The first function of the catcher is to demodulate the carrier and apply forward error correction to fix transmission errors. It extracts the IP data from the MPEG transport stream and analyzes the file to

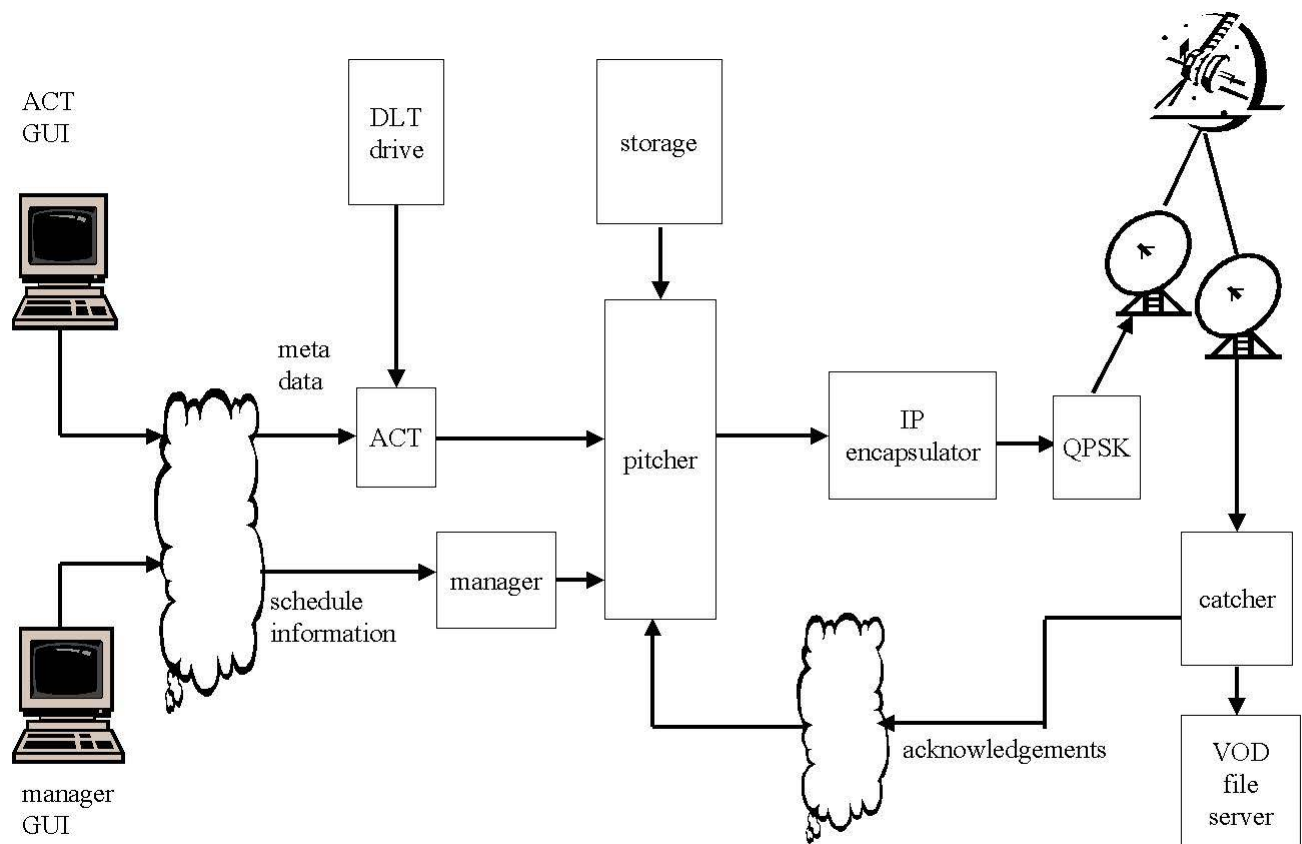
verify the transmission. If any packets are

missing or corrupted, a request for retransmission is sent over the Internet to the pitcher. The pitcher collects these requests from all the catchers and resends any corrupted packets. Once a catcher verifies receipt of a file, it sends an acknowledgement to the pitcher and the transfer to the catcher is complete.

The catcher then decrypts the content and passes it over a 100 baseT network to the VOD file server, where it is ready to be played once the movie availability date is reached.

Conclusion

A significant way to minimize the costs of operating a video-on-demand system is to centralize the encoding and distribution functions. iN Demand has developed such a



system by using an open encoding specification that supports delivery of content through a variety of file servers and digital settop vendors. Adding meta data to the movie saves labor costs at cable systems and reduces the chance for errors. Finally, an IP multicast system with triple DES encryption is used for secure delivery of content. A dedicated satellite carrier is used for VOD, allowing for a scalable, cost efficient system for transmission.

The use of a common encoding specification and distribution system allows VOD to become a viable additional revenue stream for cable operators.

HI PHY LITE – A PRAGMATIC APPROACH TO ADVANCED PHY

Dr. Ofir Shalvi, Noam Geri, Ariel Yagil, Daniel Wajcer
Texas Instruments, Broadband Communications Group

Abstract

This paper presents a new generation of Cable Modem and Cable Modem Termination Systems (CMTS) which utilize an advanced upstream Physical Layer technology known as 'HI PHY Lite' or 'Advanced TDMA', which increases the upstream capacity and improves robustness to common channel impairments while maintaining full backward compatibility to DOCSIS 1.1 systems.

In this paper we will provide an overview on the HI PHY Lite Modulation scheme, and will present data on the performance of the HI PHY Lite system in various channel. We will demonstrate how HI PHY Lite systems provide MSOs an evolutionary path to increasing the capacity of the CATV reverse channel, thereby allowing operators to offer their subscribers greater upstream bandwidth for various services such as video conferencing and telephony.

INTRODUCTION

Cable operators have been remarkably successful in establishing DOCSIS as the standard for data over cable service in North America. Today, dozens of cable modem vendors offer interoperable, DOCSIS standard-based cable modems, using chip-sets from multiple silicon vendors. The vast majority of cable modems deployed today in North America are DOCSIS based, and DOCSIS and its derivative Euro-DOCSIS are also being widely deployed in Europe and the rest of the world.

What contributed to the rapid turn around from conception to deployment of DOCSIS systems was the fact that when designing the

DOCSIS specification operators preferred mature and proven technologies over state-of-the-art. Faced with competition from DSL service providers, cable MSOs focused on getting a solid solution as fast-to-market as possible, with improvement to be considered for future generations.

Having successfully completed the first phase of establishing a cable modem standard, operators can now focus on improvements. One area for improvement is the transmission scheme used for the upstream channel. Current DOCSIS systems do not efficiently use the scarce bandwidth available to operators in many of the upstream channels. This inefficiency does not degrade service today given that penetration of cable data service is relatively low and common cable data service usage is asymmetric in its consumption of bandwidth. The most common use of cable data service, Internet browsing, is highly asymmetric in nature, requiring very little data to be transmitted in the upstream in comparison to the amount of data transmitted in the downstream. However, as cable modem penetration increases and as more symmetric services and applications such as voice over IP, video conferencing and peer-to-peer networking gain popularity, the inefficient use of the upstream channel could potentially limit operators in expanding their services. Furthermore, the shortage of bandwidth in the upstream is exacerbated by the fact that the upstream channel is plagued with various interferences, which further limit the use of the already limited upstream channel.

To address the problem of the upstream channel the Society of Cable Telecommunications Engineers (SCTE) requested from the IEEE to evaluate new

technologies for possible enhancement of the current upstream Physical Layer (PHY) specification. IEEE 802.14a working group was set to up to do this work, and to design a new PHY specification which would provide higher capacity in the upstream while at the same time providing greater immunity to common channel impairments. Although this effort, which will be described in greater detail in this paper, did not result in a completed specification, it did lead to draft specification, a derivative of which is the *HI PHY Lite* (a.k.a *Advanced TDMA*) proposal which is the subject of this paper.

THE CATV UPSTREAM CHANNEL

The upstream channel in the cable TV network covers the 5-42MHz frequency band. Due to the 'tree and branch' topology of the CATV network, the upstream is characterized by noise which is the accumulation of noise generated throughout the network. This includes the following noise sources:

- 1) White noise generated by active components on the network
- 2) Narrowband Ingress, typically originating from other transmitters such as amateur AM radio.
- 3) High rate impulse noise originating from electric current, consisting of short, typically one microsecond or shorter, noise impulses at a repetition rate varying from a few hundred occurrences per second to several thousands per second.
- 4) Low rate, but long duration, wideband bursts which could occur as frequently as every 10-20 seconds and could last 10-50 microsecond.

On top of these noise sources, the upstream signal is subject to multi-path reflections due to impedance mismatch of the plant's components and unterminated cables

PROPOSED SOLUTIONS FOR THE UPSTREAM CHANNEL AND IEEE 802.14 ACTIVITY

Three proposals were submitted to the IEEE 802.14 for consideration as next generation CATV upstream PHY (a.k.a HI PHY). The proposals included Frequency Agile TDMA (FA-TDMA), Variable Constellation Multi-tone (VCMT) and Synchronous CDMA (S-CDMA). IEEE 802.14a HI PHY working group adopted a combination of the advanced TDMA and S-CDMA (as an optional mode) proposals as the basis for its specification. For a more detailed description of the proposals and the debate in IEEE 802.14 see [1].

After a long process which spanned over a year, IEEE 801.14a, HI PHY working group was unable to reach an agreed upon specification and was ultimately disband without completing its goal.

HI PHY LITE

The *HI PHY Lite* modulation scheme was jointly developed by Texas Instruments and Broadcom Corporation [2]. It was derived from the more complicated modulation scheme proposed by the IEEE802.14a work group. However, *HI PHY Lite* excludes many of the IEEE802.14a modulation tools, such as TCM and BICM line codes, inner interleaving, non-linear precoding, CDMA modulations, and synchronous modulations, hence the use of the term 'Lite'. In our opinion, all those tools provide a relatively marginal performance gain, which does not justify the significant increase in cost, complexity and in the time it would take to the industry to reach fully tested and interoperable and multi vendor products.

HI PHY Lite modulation scheme extends the Physical layer of DOCSIS 1.0 in the following manner:

8, 32, and 64 QAM constellations

This allows increasing the spectral efficiency by up to 50% in cable plants that have good SNR, compared to DOCSIS 1.0 or 1.1 whose highest constellation is 16QAM. Such cable plants are expected to become more and more common as operators improve their plants and pull fiber deeper into the HFC.

By adding 8-QAM and 32-QAM constellations, the *HI PHY Lite* modulation scheme enables finer matching the data rate in the channel to the SNR conditions in the plant.

Symbol rate of 5.12 MHz

This allows reducing the number of headend receivers per CMTS by a factor of two, and thus offering smaller CMTS equipment, in comparison to DOCSIS 1.0/1.1, where the highest symbol rate is 2.56 MHz. It also allows improving the network efficiency of the upstream plant, as the bandwidth demand of the cable plant can be shared across a smaller number of channels.

Improved Error Correction Code

The *HI PHY Lite* Forward Error Correction (FEC) scheme includes dynamic interleaving and extends the maximum error protection ability of DOCSIS 1.0's Reed-Solomon FEC from 10 byte errors to 16 byte errors. This allows significantly higher robustness to burst noise and impulse noise.

HI PHY LITE PERFORMANCE

In this paragraph we describe the performance of a *HI PHY Lite* system which uses the new generation of Texas Instruments Headend Burst Receiver IC (TCNET4522) and HI PHY Lite DOCSIS CM IC (TCNET4042) which implement Hi PHY Lite, and utilize Texas

Instruments' INCA™ technology for ingress cancellation and capacity increase

Performance in White Noise

Table 2 shows white noise performance that can be achieved by the TCNET4522 and TCNET4042¹

Table 1: HI-PHY Lite Throughput

Constellation	Throughput [bits/symbol]	SNR @ post FEC BER=10 ⁻⁸
QPSK	1.75	8.5 dB
8QAM	2.73	13.1 dB
16QAM	3.64	16.0 dB
32QAM	4.55	19.3 dB
64QAM	5.45	22.2 dB

The performance in QPSK and 16QAM modulation is very similar to the attainable performance when using the DOCSIS scheme, beside a fraction of dB improvement which can be obtained by using the maximal T=16 error protection capability of the Reed-Solomon FEC. We observe that the 64QAM modulation, which allows 50% throughput increase compared to DOCSIS, is feasible when the SNR is 22 dB or higher. We believe that such SNR levels are becoming more and more common as cable operators increase the fiber portion of HFC plants, and clean the upstream plants. On the other hand, some of the cable plants are still not "clean" yet, and there is also a need to operate at very low SNR levels. Figure 1 depicts the ability of the TCNET4522 and TCNET4042 to operate at very low SNR levels². As we can see, one needs to sacrifice data rate in order to operate at low SNR levels, however, operation is

¹ When calculating SNR, the noise power is measured over the bandwidth of the signal (1.25 times the symbol rate)

² The noise power is measured over the bandwidth of a 5.12 Mega-symbols per second

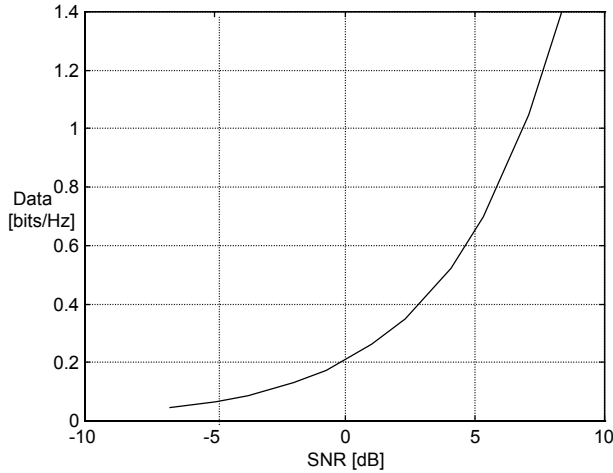


Figure 1: Operation at the Low SNR

feasible even at SNR levels of -7 dB (that is the noise power is higher by 7 dB than the signal power). Operation in these low SNR levels is achieved by retreating to the low symbol rates of DOCSIS, while the CM still maintains the high power level of a full-band 5.12 Mega-symbols-per-second HI PHY Lite signal.

Ingress Noise Performance

Perhaps the biggest advantage of the TCNET4522 burst receiver is its robustness to narrowband ingress, using the INCA[3] algorithms. Table 3 shows Carrier to Interference (C/I) that can be tolerated by the 4522 when the noise is a sine wave (CW) signal, with a BER of 10^{-8} .

Table 2: CW Ingress Noise Performance

Constellation	C/I with 4522	C/I with conventional receivers
QPSK	0 dB	9 dB
16-QAM	3 dB	16 dB
64-QAM	6 dB	22 dB

As one can see, the INCA algorithms allow improvements of 9-16 dB in the robustness to ingress noise.

The TCNET4522 and TCNET4042 were also tested with ingress noise signals recorded in noisy cable plants. Figure 2 shows the noise spectrum of a plant with heavy ingress noise. It also shows the spectrum of the signal plus noise in the case where the system operated in 16 QAM. The signal level that was required for proper operation was 15.5 dB above the power level of the noise, which is 4-5 dB better than the power level required by conventional burst receivers. We note that this figure shows the time-average of the noise spectrum, but actually the noise spectrum was dynamically time varying.

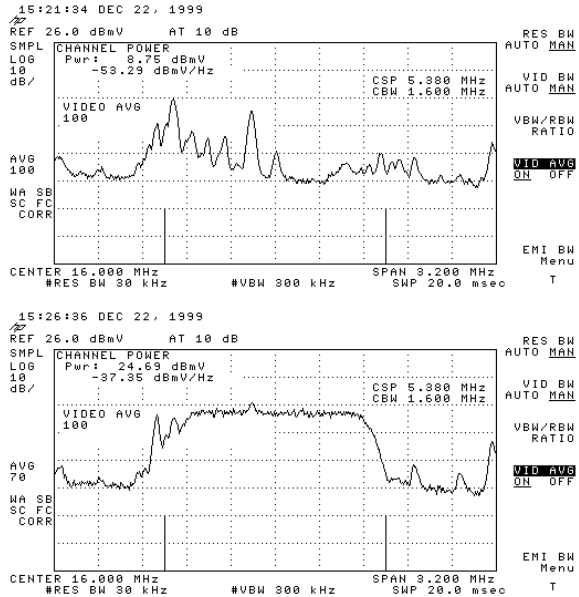


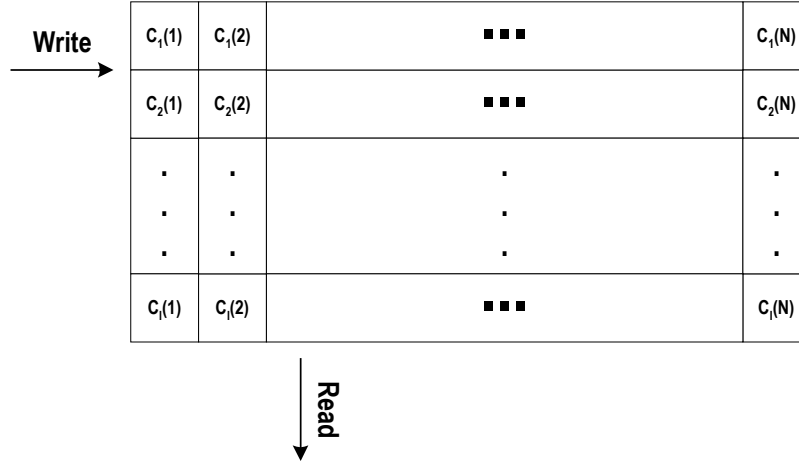
Figure 2

Burst Noise Performance

The DOCSIS 1.0 Upstream PHY combats burst noise by using Reed-Solomon error correction codes. These codes can correct up to T byte errors, where $T=0\ldots 10$. The maximal burst length that can be corrected by a DOCSIS 1.0 system is:

$$T_{\text{BURST}} = T/R_{\text{byte}}$$

Figure 3: The HI PHY Lite Interleaver



Where T is the correction factor of the code, and R_{byte} is the byte rate of the channel. Thus, for example, when the channel operates at its maximum byte rate, $R_{\text{byte}} = 1.28 \cdot 10^6$ bytes per second, and its maximum error protection capability, $T=10$, the maximal burst length that can be tolerated is 7-8 microseconds.

The *HI-PHY Lite* scheme employs a 2048 bytes block byte interleaver that allows a significantly improving burst correction capability. The byte interleaver is shown in Figure 3. It is a two-dimensional memory structure with N columns and I rows, where N is the Reed-Solomon code-word length. The interleaver performs permutations on a block of $N \times I$ data bytes in the following way: The bytes are written into the interleaver row by row, and read from the interleaver column by

column (see Figure 3). The CM and CMTS dynamically program the interleaver's parameters in order to optimize them to the size of each data packet.

When a byte interleaver is used, the maximum burst length that can be corrected by a *HI-PHY Lite* system is multiplied by the interleaver depth (I), that is:

$$T_{\text{BURST}} = I \times T / R_{\text{byte}}$$

When the SNR in the channel is high, higher constellations can be used in conjunction with low code-rate Reed-Solomon code and byte interleaving to further increase burst noise robustness.

Figure 4 shows the attainable throughput of a *HI-PHY Lite* system as a function of the burst

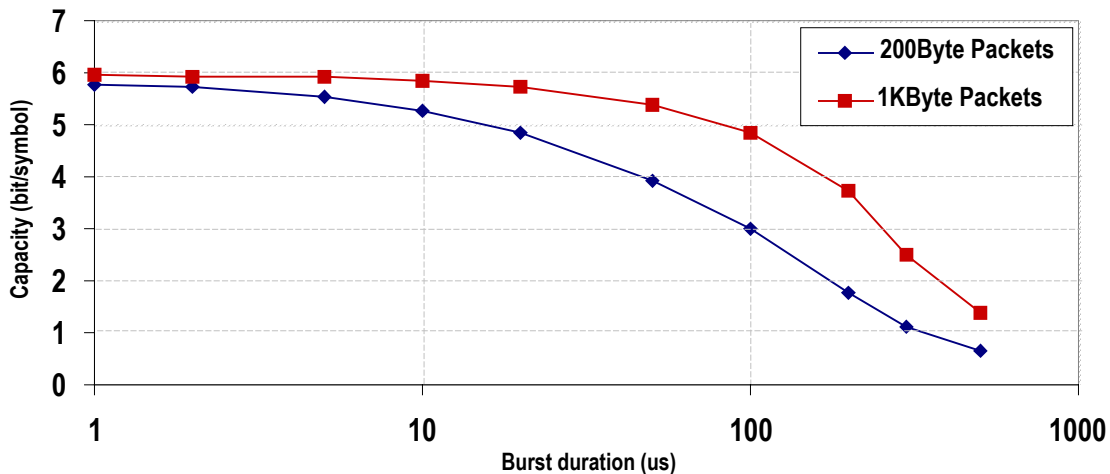


Figure 4. HI PHY Lite Burst Tolerance ($R_s=1.28\text{Mbaud}$)

length that the system can tolerate.

Performance is analyzed for 200 byte packets and 1Kbytes packets. In this test we have used error correction factor $T=10$. An improvement by up to a factor of 1.6 in burst duration can be obtained by using the maximum error correction factor of *HI PHY Lite* FEC, $T=16$.

For the 200 byte packets, interleaving is limited, thus burst tolerance performance is degraded. In comparison, DOCSIS 1.0 frames, which do not have any interleaving, can tolerate bursts that are shorter than 30us and 13us for QPSK and 16-QAM respectively. The longest burst that could be tolerated in 64-QAM without interleaving is 7u.

Note that when extremely large noise bursts are present, they can be tolerated by reducing the baud rate (burst tolerance, T_{BURST} , is increased by the factor of baud rate decrease through the decrease in R_{byte}).

We observe that the *HI-PHY Lite* system can easily tolerate bursts of up to 10 microseconds, which covers the vast majority of the noise bursts in the system. By gradually trading throughput for burst noise robustness, the system can tolerate bursts of up to 500 microseconds, which are considered to be very rare.

It is important to note that the burst robustness is proportional to the symbol rate, and it can be further improved by up to a factor of 8 by using the lower symbol rates of DOCSIS. This increases the number of upstream channels and may sacrifice MAC layer efficiency, but it provides a pragmatic solution for the rare cases of frequent occurrence of very long noise bursts (more than 100 microseconds).

Finally, we note that the TCNET4522 employs signal processing algorithms that detects noise bursts and allows improvement

of the maximum burst length protection by up to a factor of two. In this case:

$$T_{BURST} = k_{BURST} \times I \times T/R_{byte}$$

Where k_{BURST} is in the range of [1,2], depending on the characteristics of the signal (e.g. its constellation) and the noise burst (e.g. its power level).

We conclude that by using *HI-PHY Lite*, the cable operator can achieve extremely high robustness to long noise bursts by using low symbol rates and/or slightly lower throughput.

Impulse Noise Performance

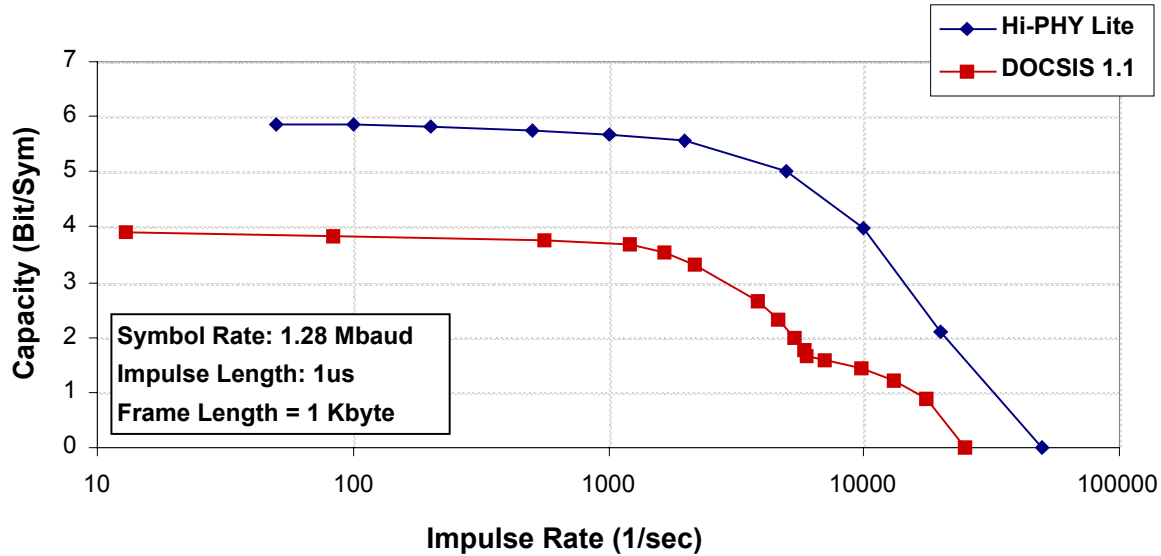
Figure 5 shows the attainable throughput of the DOCSIS 1.0 and *HI-PHY Lite* in a 1.28M symbols per second channel with 1000 bytes data packets. The noise impulses are 1 microsecond each, and the SNR is -5dB within the impulses and 25dB otherwise. The impulses appear randomly in time (Poisson distribution) and the horizontal axis of the figure is the average number of impulses per second. In this tests we have only used FEC error protection factors of $T=0...10$, and some improvement in performance can further be gained by using $T=11...16$.

We observe that the HI-PHY Lite modulation scheme can easily tolerate 1000 impulses per second and by trading data rate it can tolerate up to 25000 impulses per second. The DOCSIS 1.0 modulation scheme is inferior to HI-PHY Lite mainly due to lacking higher constellations. However, both HI-PHY Lite and DOCSIS 1.0 are capable of operating in scenarios of 1000-10000 impulses per second which are believed to represent the toughest upstream channels.

Inter-symbol-Interference

As in DOCSIS 1.1, the CM and CMTS support a T-spaced linear pre-equalizer in order to address reflection in the upstream

Figure 5: Hi-Phy Lite - Capacity vs. Impulse Rate



channel. While in DOCSIS 1.1 the pre-equalizer has 8 taps, in the new *HI PHY Lite* modulation scheme, the pre-equalizer was extended to 24 taps to support the larger symbol rate and the higher constellations.

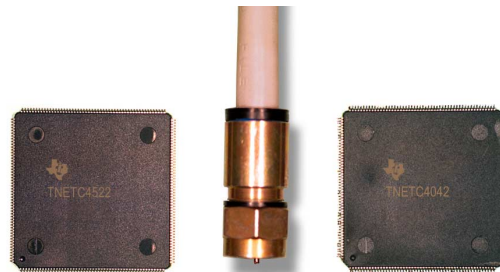
SILICON IMPLEMENTATION

Figure 6 shows a picture of the TCNET4522 and the TCNET4042 chips. The TNETC4522 is a dual burst receiver that supports *two HI PHY Lite* upstream channels of up to 30 Mbps each. Its total throughput reaches up to 60 Mbps, which translates to a significant improvement in the throughput per board area, compared to previous solutions that supported a single DOCSIS channel at up to 10 Mbps. The TNETC4522 directly samples the whole IF spectrum and the required upstream frequencies are digitally demodulated. This scheme eliminates the necessity for a tuner.

The TNETC4042 is an integrated cable modem chip, which fully supports the *HI PHY Lite* modulation scheme. The chip is backward compatible with the DOCSIS 1.0 and DOCSIS 1.1 standards, and a modem

built around TNETC4042 has gained CableLabs DOCSIS1.0 certification.

Figure 6



SUMMARY

We have presented a new generation of CM and CMTS IC's which implement HI PHY Lite modulation scheme and incorporate INCA™ technology. The attainable performance of these IC's was demonstrated in this paper.

The *HI PHY Lite* specification emphasizes simplicity, multi-vendor interoperability and backward compatibility. While

communication technology does offer building blocks such as TCM, Precoding etc, which could possibly enhance the performance of the system, it is our opinion that the performance gained by these enhancements would be marginal, not justifying the added complexity which would ultimately impact multi-vendor interoperability and cost.

This new technology enables significant improvement to upstream throughput allowing operators to offer subscribers services which demand greater bandwidth in the upstream such as VoIP and Video conferencing.

REFERENCES

- [1] “Advanced Modulation Schemes For Cable TV Upstream Channel”, Ofir Shalvi and Noam Geri, ICCE Conference Proceedings, June 2000
- [2] “Advanced TDMA Proposal for HFC Upstream Transmission”, Ofir Shalvi, Jonathan Min et al.(Texas Instruments- Cable Broadband Communications and Broadcom Corp.), Oct 1999
- [3] INCA Technology White Paper (Texas Instruments –CBC),
www.ti.com/sc/docs/innovate/cable

Home Networking: Extending Cable's Reach

Samuel H. Russ
Subscriber Networks
Scientific-Atlanta

Abstract

Home networking offers the promise of distributing high-speed Internet access and entertainment media throughout the home. By pushing the Internet experience out to where non-technical users are, and by making the in-home cable system the anchor of a pervasive communications and entertainment experience, it can increase penetration and reduce churn.

Scientific-Atlanta has experimented with three different wireless standards and a variety of devices in order to understand how home networking can become a reality. Our experiments have highlighted several short-term obstacles that must be overcome as home networking becomes a routine and lucrative addition to an MSO's product offering.

Survey of Applications and Physical Media

What's It Good for?

Home networking is rapidly emerging as a key technology, a rapidly growing market, and a prime opportunity for forward-thinking companies. There are (at least) four distinct roles for home networks—routing video, routing voice (presumably over IP), routing data, and basic home automation functions. This is important because they have completely different bandwidth needs and realtime / quality of service constraints, as shown in the table below.

This chart indicates that several different uses for home networking exist, with widely disparate service requirements. A comprehensive home networking solution must take these factors into account. For example, power line networking may be widely adopted for home automation functions (such as climate control, appliance access, and home security) and phone line networking for data transfer.

Application	Bandwidth needed	QoS Guarantees?	Two-way?
Video: HDTV	19 Mbps	Yes	No
Video: DVD	3 - 9 Mbps	Yes	No
Video: MPEG-2	2 - 5 Mbps	Yes	No
Video: Videoconf.	30 kbps - 1.5 Mbps	Yes	Yes
Data: File Sharing	200 kbps – 10 Mbps	No	Yes
Data: Printing	60 kbps – 1 Mbps	No	No
Data: Web Surfing	80 – 250 kbps	No	No
Interactive Gaming	10-100 kbps	Yes	Yes
Audio: "Hi-Fi"	100 – 200 kbps	Yes	No
Audio: Voice	5 - 64 kbps	Yes	Yes
Home Automation	100's of bps	No	No

Table 1: Some Applications for Home Networking

One interesting point of this comparison is the “1 Mbps gap”. Essentially all applications except video work well, or at least are usable, below 1 Mbps. Broadcast-quality video, on the other hand, demands more than 1 Mbps. This highlights the importance of understanding what services subscribers will want from home networking as operators plan for the subscriber products and the infrastructure required to support home networking.

How Can I Connect?

There are seven ways to convey voice, video, and/or data inside a home. One, ultrasonic, will not be considered here. The remaining six are important and should be considered in an overall home networking strategy.

The six media are radio, infrared, phone line, power line, cable coax, and wired networking, as illustrated below in Figure 1.

Some of these media encounter serious limitations immediately. For example, wired networking (Ethernet and Firewire) require extensive rewiring (in the case of Ethernet) or only work over short distances (in the case of Firewire). Infrared is limited as a practical matter to line-of-sight and is really only possible inside a single room.

Cable coax can form a home networking backbone, and well-established standards exist for it, but it is limited by the number of outlets inside a house. Its role in home networking will grow in the future, but since it is so well-understood by this audience it will not be considered further.

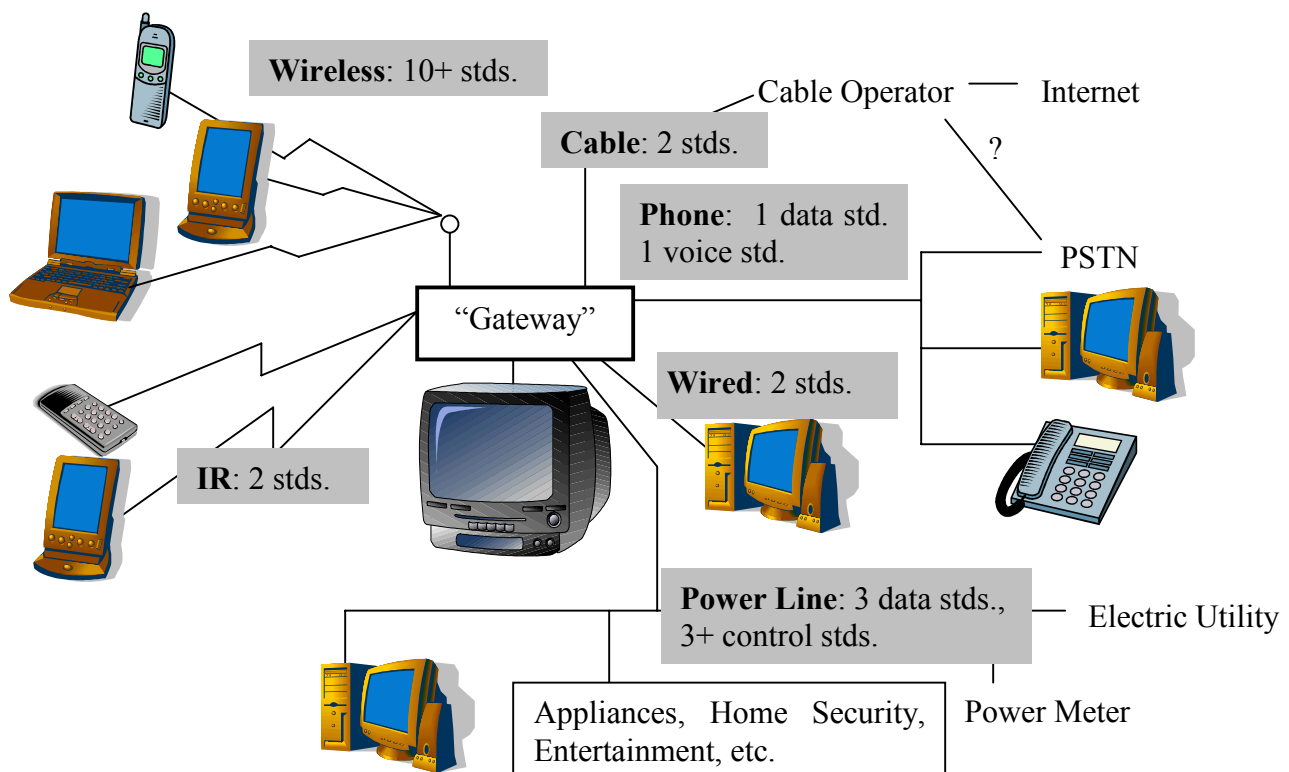


Figure 1: Home Networking Physical Media

There remain, then, three additional physical media to consider in developing a home networking strategy -- RF, phone line, and power line.

RF networking (or “wireless networking”, which is a bit of a misnomer since IR is also wireless) is emerging as a powerful and flexible way to connect devices inside a home. It uses no new wiring and offers portability and ubiquity advantages over wired networking. However, it is not as foolproof as wired networking, and the market has been hindered by profound standards confusion.

There are three standards available today, Bluetooth, IEEE 802.11b (hereafter referred to as “802.11b”), and HomeRF, that are market leaders. Their capabilities will be discussed below.

Some of these wireless standards have interference problems with each other, although this is becoming less of an issue. Two of the major standards, Bluetooth and 802.11b, are being harmonized under IEEE 802.15.2, which may help solve some interference issues. Another source of reconciliation is a move towards 5 GHz networking, and many consider the way of the future to be 5 GHz 802.11a and 2.4 GHz Bluetooth.

Phone line networking takes full advantage of the inherent signal quality of twisted pair. It can routinely deliver tens of megabits per second (Mbps) around typical homes, making it a serious candidate for essentially all home networking applications. More importantly, there is a single, uniform, universally adopted standard – HomePNA. However, phone taps are limited, especially in Europe, and it always seems like set-top boxes are never near RJ-11 jacks.

Power line networking has the advantage of a truly ubiquitous network interface, and so is as convenient as wireless (unless you need portability). There is still confusion over the standards, and it is not even clear that all of the standards will work in the field. Proposed data rates are as high as 14 Mbps, but the usable bandwidth must be divided by the number of homes that share a transformer’s secondary winding. (In the US this is typically 5-10 homes.) One concludes, then, that powerline is a serious candidate for applications requiring 1 Mbps or less.

Commonly discussed multi-Mbps powerline standards include HomePlug and Adaptive Networks’ solution. Both are in prototyping and field-trial stages.

Our Experiments

Scientific-Atlanta assembled a series of experiments in the area of wireless (RF) home networking in order to demonstrate capabilities and to learn first-hand the requirements for a complete product offering.

First, we connected an Intel AnyPoint HomeRF-based access point through an Explorer set-top box’s USB port. This enabled the delivery of 1.6 Mbps of TCP/IP to a remote portable client. The portable client, in this case a laptop, was able to browse the web and watch low bitrate movie trailers. These trailers typically use around 300 kbps and are CIF or QCIF resolution screens. The sustained throughput of HomeRF is actually about 550 kbps; 1.6 Mbps is the physical layer symbol rate and neglects effects such as error correction and packet header overhead.

The system was decidedly popular at a cable trade show. One non-technical user even commented that while he currently lived outside of one of Scientific-Atlanta’s customer’s coverage areas, he would move into their service area if they launched wireless web access!

Second, we connected a set-top to a Bluetooth-enabled web tablet made by Ericsson. Bluetooth was routed into the set-top through an external attachment for demonstration purposes. The web tablet was able to display a program guide, and selections made on the web pad changed the set-top's channel. This demonstrates the ability of wireless peripherals to operate two-way and to function as a remote control. Note that the Ericsson product was also able to browse the web, although Bluetooth provides a relatively lower 721 kbps bitrate. This sort of remote control application is exactly the type of short-range, low-bitrate application that the designers of Bluetooth had in mind.

Third, with help from Intersil, we put an 802.11b access point inside an Explorer 8000 set-top. We encoded the set-top's tuner

output using 1.5 Mbps MPEG-2 and transmitted it via the access point to a remote laptop. The remote laptop simultaneously displayed a live web browser, a functioning MP-3 player, a live video screen, and a working remote control. The remote control tuned the set-top, activated the program guide, and can even work with interactive set-top applications like pizza on demand. A screen capture of the demo is shown below in Figure 2.

The three demonstrations were actually laid out side-by-side at the trade show (against the advice of engineers) and actually interoperated fairly well. The three displays had to be kept near their respective access points, and carrying one display next to another access point did disrupt the traffic, but it did highlight that non-interfering operation may not be as difficult as some fear.



Figure 2: Screen Capture of Scientific-Atlanta / Intersil 802.11b Demonstration

Capabilities of Existing Systems: How They Can Help MSO's

These demonstrations highlight four key advantages of home networking.

First, basic web access can be pushed out to where non-technical users are. This is an extremely important paradigm shift. Think about technophobes like my wife. They hate using the Internet because the computer is off in a back room and, if they rely on dial-up, there is a lengthy and annoying process to get connected. Wireless networking through a DOCSIS channel changes all that because wireless is available anywhere and DOCSIS is always-on. This is one strength of DOCSIS-based home networking that may increase penetration rates of digital cable.

Second, subscribers purchase the peripherals that hang off of the home wireless network, such as extra PC and PDA network adapters and wireless appliances like MP3 players. Once they have purchased the peripherals, they will be strongly motivated to keep their existing service. Thus home networking is a powerful churn-reduction tool.

Third, deploying the infrastructure enables lucrative services. For example, the MP3 player we showed in the Intersil 802.11b demonstration foreshadows possible new subscription services.

Fourth, home networking is on the verge of being able to deliver broadcast-quality video in the home. The existence of the cable industry is proof that video offers tangible revenue streams, and this is emerging as a new way to distribute video content.

Limitations of Existing Systems: Lessons Learned

These demonstrations raised several key issues related to the “productization” of home networking.

Fire Hose or Soda Straw?

The first issue that comes up is bandwidth. As seen in the first table, subscribers' bandwidth needs are highly application-dependent. Some networking standards are not suitable for all applications.

For example, phone line is currently the only technology with ample bandwidth for video, and 802.11b is the only open standard in the wireless world that has barely enough bandwidth for video. (There is one proprietary system, ShareWave WhiteCap, which increases the usable bandwidth available for video delivery).

The answer to the question “how much is enough?” is simply “it depends on what you want to do with it.” In the video-centric world of cable, however, home networking is really only barely able to deliver video right now.

The IEEE 802.11a standard is on the horizon; most expect it to be generally available by the end of 2001. It features a substantial increase in bandwidth that, coupled with the quality of service (QoS) features added by IEEE 802.11e, may well enable practical multi-channel wireless video delivery in the home.

Hey, I'm Trying to Talk!

The 802.11b demo highlighted the need for quality of service in delivering video. The 1.5 Mbps MPEG-2 was selected simply to leave enough bandwidth margin for web browsing and MP-3 in the absence of QoS. (We have to be honest here – 1.5 Mbps is not broadcast quality, which is why this was a demo and not a real product.)

Note that “quality of service” in the context of broadcast-quality video implies not only reservation of bandwidth but also guarantees and bounds on packet jitter. Web browsing over the corporate LAN, which boosts the peak bandwidth of a web page download, can

disrupt the demo's video panel, for example, because it does not leave enough free bandwidth for video. Likewise, attempting to display a movie trailer in the demo web browser panel completely wipes out the demo video panel.

By the way, in the case of 802.11b this is being handled two ways. First, the IEEE 802.11e standard is being constructed to reserve bandwidth on an 802.11b channel. Second, proprietary standards such as ShareWave WhiteCap are also available.

Perform Steps 1-37 in Painstaking Detail

Finally, we learned through the process of nailing up trade show demonstrations that installing, maintaining, and administering home networks are going to be hard. Do set-tops have to run DHCP? What about DNS and NAT? Do they run Firewalls and, if so, how much more expensive is it to run firewall (because of increased memory and processor demands)? What happens when there is interference on the currently selected 802.11b channel? Who decides there needs to be a

channel-change? How does the operator know the state of the 802.11b access point, and how does the operator change the settings? Thankfully these and a myriad of other small but important details are being taken up in the CableLabs' CableHome home networking forum.

Where We Need to Go: Conclusions and Future Work

By exploring various technologies and various product form factors, it became clear that home networking is nearing fruition, but not quite ready yet. Various standards-making efforts, including those supporting systematic support of home networking (CableHome) and those adding quality of service to Ethernet-based protocols (IEEE 802.11e) can close the gap. Practical home networking will soon be in a position to be a valuable addition to MSO product offerings.

Acknowledgements

The author wishes to thank Tony Wasilewski and Bill Wall for their helpful review of and comments about this paper.

HOP A RIDE ON THE METRO GETTING YOUR IP PACKETS FROM HERE TO THERE

Erik Metz and David Grubb
Motorola Broadband Communications Sector
John Brenneman and Daniel Sills
Geyser Networks Inc.

ABSTRACT

Next generation Synchronous Optical Networking equipment is paired with data centric interfaces and packet switching engines to provide reliable Metropolitan networking for Headend interconnects. This integration has lead to new developments in multiplexing structures and bandwidth allocations. By creating more granular packet containers, and dynamic allocation, services are more efficiently packaged onto the transport network. Implementing priority based queuing allows for services to be expedited onto the network. Services bound by Quality of Service constraints in delay and

delay variance can be individually managed across the transport network. Differentiated Services, Open Access and Carrier Class transmission will merge into one consolidated backbone serving entire metropolitan regions.

WHAT SEEMS TO BE THE PROBLEM?

With the increased deployments in cable modem services, and the looming explosion of Voice over Internet Protocol (VOIP) telephony services, getting all of the packets to their destinations on time will play a critical role in the Metropolitan Network. Multiple System Operators (MSO) have many interconnected Headends serving geographical

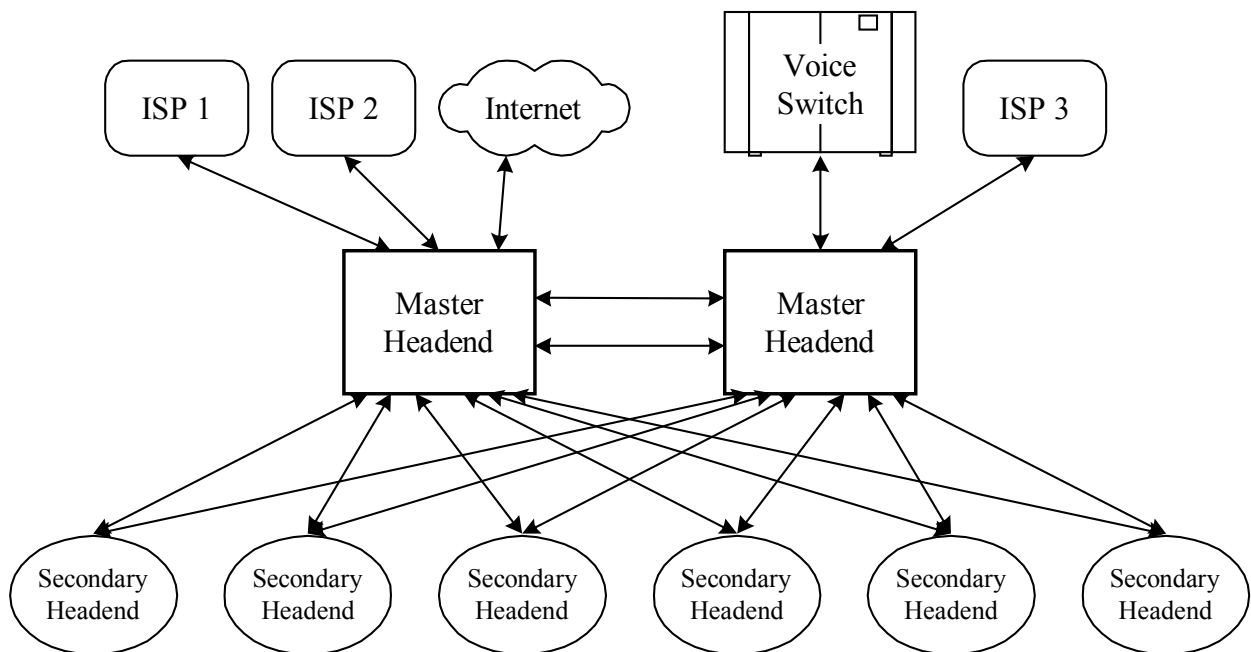


Figure 1

regions that make up the metropolitan network. One or more of these Headends may provide access to the Internet and one or more of these may have access to a voice switch for telephony. These Headends act as gateways to their respective services. These will be called the Master Headends. The remaining Headends will be called Secondary Headends. The Secondary Headends terminate the upstream and downstream traffic to/from the Hybrid Fiber COAX (HFC) plant via the Cable Modem Termination Shelf (CMTS) devices. The input and output from the CMTS to the Metropolitan Network is Internet Protocol (IP) packets. These packets need to be separated based on the type of service they are carrying and by destination Internet Service Provider (ISP).

The traffic that flows in this type of arrangement can be categorized as “Hub and Spoke” traffic (Figure 1). This implies that there is a large amount of traffic flow into and out of the Master Headends, and less in each of the Secondary Headends. This is a logical star. If there are two Master Headends, then it will look like two overlapping stars. Each Secondary Headend will be generating traffic that is different for each type of service it is providing. Internet traffic is bursty when looked at over short periods of time and when aggregated, can have sustained flows that change dramatically over the period of a single day. VOIP traffic is not necessarily bursty, but changes deterministically with each phone call that is placed. These aggregated traffic patterns will also change dramatically throughout the period of a single day.

Clearly there needs to be a solution that can efficiently aggregate the traffic from the Secondary Headends and reliably connect the Master Headends for access to the necessary gateways.

NEXT GENERATION SONET/SDH

Next generation Synchronous Optical NETworking (SONET)/ Synchronous Digital Hierarchy (SDH) platforms incorporate the latest technologies in data networking with the reliable transport from the telecommunications arena. SONET/SDH has been the benchmark for reliability for the Regional and Long Haul telecommunication networks. It is based on self healing rings and highly available equipment shelves with sub 50 ms recovery times from catastrophic failures such as a fiber cut. It is standardized world wide and has been deployed since the 1980's.

Traditional SONET/SDH was developed with the transport of multiplexed voice signals in mind. The basic building block is a Synchronous Transport Signal 1 (STS-1), or 55 Mbps, or Synchronous Transport Multiplex 1 (STM-1) for SDH, or 155 Mbps containers. SONET/SDH Networks are very scalable, and are available with different line rate interfaces and different trunk rate outputs (Table 1). Many of today's SONET networks are OC-48, (Optical Carrier 48 or 48 STS-1's) with a line rate of 2.455 Gbps. These large statically assigned containers make it difficult to implement today's packet services efficiently.

SONET Hierarchy		
Signal	Bit Rate	Capacity
STS-1, OC-1	51.840 Mb/s	1 DS3 or 28 DS1s
STS-3, OC-3	155.520 Mb/s	3 DS3s or 84 DS1s
STS-12, OC-12	622.080 Mb/s	12 DS3s or 336 DS1s
STS-48, OC-48	2488.320 Mb/s	48 DS3s or 1344 DS1s
STS-192, OC-192	9953.280 Mb/s	192 DS3s or 5376 DS1s
STS-768, OC-768	39,813.12 Mb/s	768 DS3s or 21,504 DS1s

Table 1

Next Generation SONET/SDH networks are implementing significantly more granular transport containers. These containers are allocated and de-allocated at the physical layer and the allocation process is managed by hardware based controllers that operate in near

real-time. This near real-time allocation allows for packet based data to be statistically multiplexed onto the transport network. That is to say, that as one service has traffic bursts, that more of the transport network can be allocated to it, over another service that may be temporarily idle. With many packet services connecting to the transport network, the hardware can provide a more statistically packed service, thus minimizing the amount of wasted bandwidth.

This process does not operate unmanaged. The bandwidth is managed on a per service basis, and is allocated using traffic descriptors. This allows differentiated services, such as VOIP, to have a significantly higher priority, thus preserving the Quality of Service (QoS) that is required for voice

traffic. The traffic descriptors can be allocated to provide a three tiered approach to the bandwidth allotment on a service by service basis. Each service will be assigned a minimum rate, that is the amount necessary to sustain the lowest amount of traffic estimated, or just enough to maintain a circuit. The next tier is the guaranteed maximum rate. This is the amount of the transport network that will be made available to the service with no statistical degradation. The last tier is the maximum best effort. This is the maximum amount of bandwidth that any service attain from the transport network, if it is available. This is also the bandwidth that must be relinquished if other services are requesting guaranteed bandwidth. This could return the service to its maximum guaranteed amount.

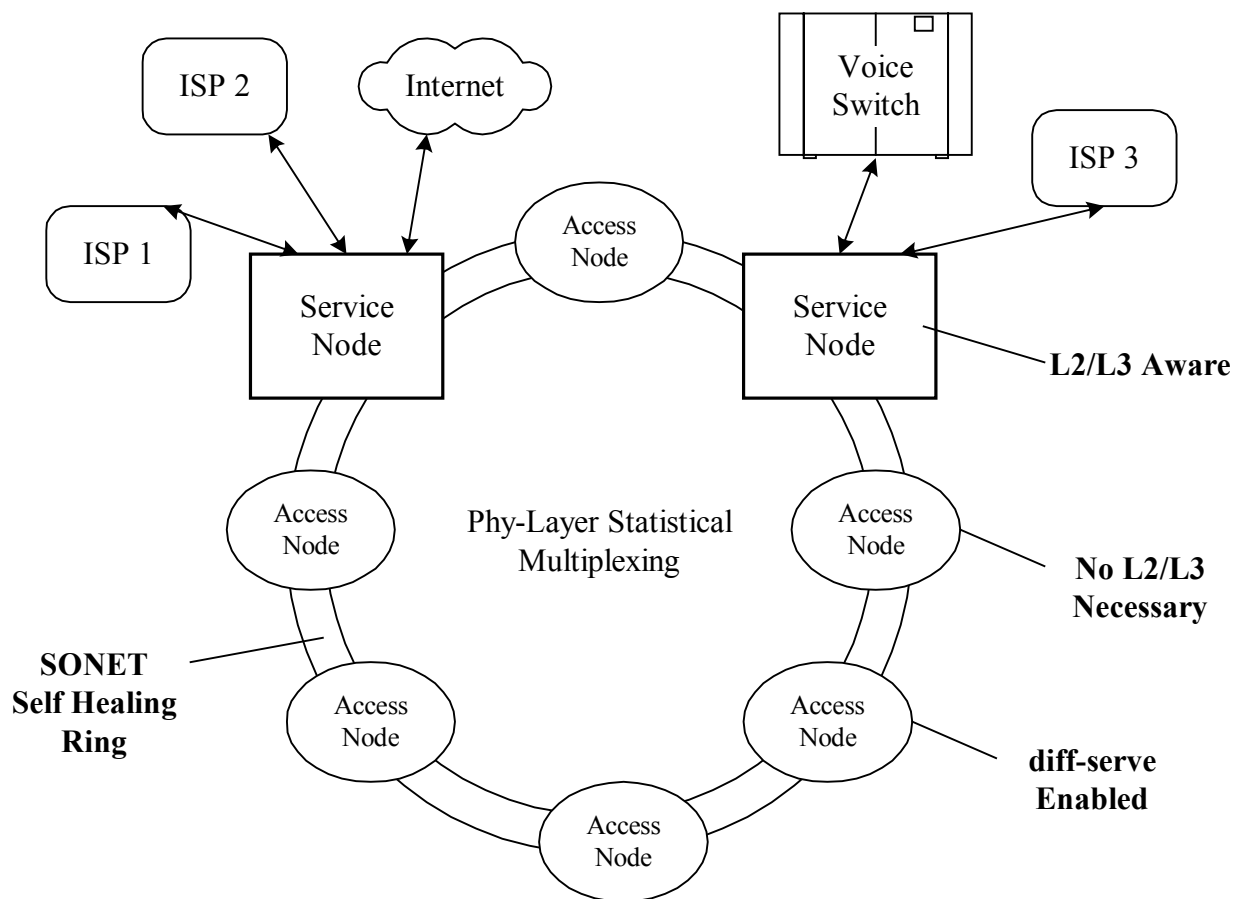


Figure 2

As an example, a 100 BaseT data service may be allocated a minimum of 10 Mbps, a guaranteed rate of 50 Mbps and a best effort rate of 100 Mbps. For a 100 BaseT VOIP service, the minimum rate may be set at 30 Mbps, a guaranteed rate of 100 Mbps, and no best effort. This would be necessary in order to never “drop” a phone call. By building these traffic descriptors on a service by service basis, the operator is able to tailor the statistical nature of the multiplexing on the transport network, and still preserve QoS.

RING SOLUTION

Using SONET/SDH rings to connect each of the Headends, the overlapping logical star becomes a single physical ring connection (Figure 2). With the Hub and Spoke service solution, the Headends become node sites on the ring. The Master Headends act as Service Nodes, and the Secondary Headends act as Access Nodes. The Access Nodes provide multi-protocol, multi-service I/O interfaces for connection with the traffic generating HFC equipment such as the CMTSs, or Host Digital Terminals (HDTs). The traffic from the Access Nodes is Classified, then policed and shaped prior to being multiplexed onto the transport network. The individual services are sent to the appropriate Service Node. The Service Node provides the necessary Layer 2 switching, or Layer 3 routing of the incoming services to the appropriate destination.

Since the Access Nodes can greatly outnumber the Services Nodes, and there is no initial need to provide L2/L3 capability at the Access Node (because most of traffic is destined for service node), the management of

this network can be significantly reduced. Furthermore, since the multiplexing is managed at the hardware level, and there is a system wide manager of such, the allocation of the bandwidth to services is performed in an unbiased nature. This also greatly simplifies the management which is already a benefactor of the existing SONET/SDH configuration policies.

The Service Nodes essentially become the brains of the Metropolitan network. They are the interface to the critical services, and provide the demarcation for multiple ISPs. If necessary, any Access Node is capable of becoming a Service Node. Also, Service Nodes can either provide unique functionality at each site, or provide replication of services from another Service Node, or simply be a hot standby.

Since the statistical multiplexing takes place at the physical layer, transmission through intermediate nodes around the ring does not compromise QoS. With input classification, prioritization, and traffic descriptors building the transport network rules and regulations, it is possible to provide effective and efficient utilization of the Metropolitan Network bandwidth. In doing so, minimizing the amount of necessary equipment and increasing the ease of management of the network, operators can capitalize on this emerging technology today, and scale it to future needs as they arise.

Erik C. Metz
Systems Engineer
Motorola BCS
101 Tournament Drive
Horsham, Pennsylvania 19044
(215) 323-2284 (phone)
(215) 323-2349 (fax)
emetz@gi.com

John Brenneman
Director of Systems Engineering
Geyser Networks
535 Del Rey Avenue
Sunnyvale, California 94085
408-331-7133 (phone)
408-830-9099 (fax)
jbrenneman@geysernetworks.com

HOW NEXT-GENERATION BROADBAND NETWORK ARCHITECTURES WILL DELIVER INTEGRATED SERVICES TO BUSINESS CUSTOMERS

Vikram Saksena
Narad Networks

Abstract

This paper describes a next-generation broadband network architecture for providing 100 Mbps/1 Gbps Ethernet services over an existing HFC infrastructure. By deploying Ethernet switching elements with integrated high speed modems at critical junction points, today's tree-and-branch cable network topology can be augmented with a symmetric, full duplex, data channel that provides orders of magnitude higher bandwidth than is currently available. With built in multi-service QoS capabilities, this next-generation broadband network becomes fully capable of addressing the large, highly lucrative commercial market segment.

Coupled with advanced network management and service delivery back-office systems, the entire process of deploying and managing broadband services can be fully automated. High levels of network reliability can be maintained by making the network elements remotely manageable and self-configurable. By radically compressing service deployment and provisioning intervals, cable companies can accelerate the process of new revenue generation. Customers get the benefit of being able to order services on demand from a broadband services portal without waiting for long service turn-up intervals and without worrying about bandwidth constraints. Flow-through provisioning and activation of back-office billing and customer care systems are accomplished through a robust Directory Enabled Networking (DEN) based service delivery platform.

INTRODUCTION

Recent years have seen the cable industry deploy a number of advanced services in the consumer marketplace. Driven by the need to generate new sources of revenue, the industry is adopting new technology and is rapidly transforming itself into one that provides multiple integrated services. This has resulted in increased cable modem penetration, expanding deployment of video-on-demand and digital TV, and cable telephony services. To accommodate these services, the bandwidth used by cable systems has steadily increased from 270 Mhz to 860 Mhz HFC over the last twenty years and the networks have been upgraded to support two-way traffic.

To continue on the revenue growth curve, the next major challenge for the cable industry is to address the needs of the commercial market. The small and medium business (SMB) segment of the commercial market in particular represents a lucrative growth opportunity for the cable companies and is a natural progression from the consumer markets addressed by the cable companies today. However, the existing cable infrastructure with its broadcast characteristics and asymmetric bandwidth capabilities is largely incapable of fully supporting the needs of the SMB market. This is supported by a recent Morgan Stanley study that reveals that the cable modem penetration in the commercial markets is insignificant when compared to DSL and other access technologies, and is not expected to increase in the foreseeable future. Recognizing these

limitations, some cable companies have deployed separate physical networks based on ATM, SONET and other proprietary access technologies to address the commercial market. However, the high cost of deploying and operating a separate physical network has limited the addressable market to a few concentrated areas that are largely beyond the reach of the majority of SMB customers both in terms of availability and costs.

This paper proposes a new broadband network architecture that leverages a cable company's largest deployed asset - the HFC network - to reach a significant portion of SMB customers and serve them in a highly efficient and cost-effective manner. The proposed solution adapts switched Ethernet technology commonly deployed in today's enterprise networks to broadband HFC networks and enables a rich suite of IP services – high speed Internet access, virtual private networking (VPN), business grade telephony, web hosting, content distribution, and storage area networking. The technology foundation, based on high speed Ethernet switching, provides fiber-like reliability and performance characteristics over existing HFC networks at a fraction of the cost of alternative fiber-based solutions. This paper also introduces a service delivery platform that enables the rapid creation and provisioning of services in near real-time using a Directory Enabled Networking (DEN) architecture.

This next generation business solution gives the cable operators a true commercial-grade broadband network and service delivery platform that transparently coexists on their ubiquitously deployed HFC network. This new architecture provides a competitive advantage against the incumbent carriers who offer conventional narrowband services in a relatively expensive and inflexible fashion due to tariff restrictions and culture.

THE SMB OPPORTUNITY

Market Size

There are approximately 8 million small and medium businesses in the US that typically employ 5 to 100 employees. Approximately 60% of these businesses are within a few hundred feet of the local HFC network and can be cost-effectively reached by the cable operator. On average an SMB pays about \$6000 in annual revenues to its service provider, which in most cases is the incumbent local exchange carrier (ILEC). Therefore, the addressable market opportunity for the cable industry is an annual \$28.8 billion. Compared to the consumer market opportunity of \$35 billion, the SMB market represents a significant growth opportunity for cable companies. Cable operators can significantly grow their revenues and EBITDA margins with modest penetration rates, given the higher margins associated with the SMB revenues.

SMB Requirements

The core set of services demanded by SMB customers includes high speed Internet access and telephony. In many cases, SMBs also require additional value-added services such as Centrex, virtual private networking, web hosting, and firewall services. Most of these services require symmetric bandwidth. A common service configuration includes an integrated access device (IAD) on the customer premises that aggregates voice and data traffic over a symmetric, full-duplex access line such as T1 or fractional T1 to the ILEC's central office (CO). At the CO, traffic is demultiplexed; voice traffic is sent to a Class 5 switch that connects to the PSTN and data traffic is sent to an IP network that provides Internet and data services.

While the 1.5 Mbps bandwidth provided by a T1 line is generally adequate for basic voice and data services, bandwidth requirements increase when customers begin demanding additional value-added services. Although NxT1 services are available from the ILECs in a few areas, by and large the next step up in access bandwidth is DS3 (45 Mbps). However, DS3 costs are prohibitively expensive and beyond the range of affordability of most SMB customers. Therefore, in many cases, as traffic demand increases customers are forced to use techniques such as voice compression to squeeze more traffic into the same T1 line. This not only leads to degradation in voice quality but also causes higher delays and packet loss for data traffic. In this scenario, lack of affordable bandwidth from the ILECs becomes a barrier to the growth and proliferation of broadband services.

It has become a common practice for the commercial customer to enter into service level agreements (SLAs) with its service provider that binds the service provider to contractual commitments related to service performance metrics such as delay, throughput, packet loss, and availability. In order to effectively deliver on these SLAs, the service provider needs a QoS-enabled, multi-service network infrastructure that is highly reliable and scalable. In addition, the service provider also requires a highly sophisticated back-office operations and management infrastructure that is fully automated to reduce provisioning cycle times and mean time to repair activities.

The incumbent service providers are at a competitive disadvantage because of their high cost legacy networks and outdated back-office operations systems that create unduly long provisioning intervals leading to customer dissatisfaction. This scenario presents a significant opportunity for the cable companies to leverage their HFC network to

offer broadband services to commercial customers at the lowest cost per bit. By adopting state-of-the-art back-office support systems, cable companies can deliver on the promise of customer SLAs in a fashion far superior to the incumbents.

CURRENT HFC PLANT CHARACTERISTICS

In current HFC networks, the coaxial portion of the plant is a tree-and-branch topology with a shared LAN-like structure. In the downstream direction, RF signals are broadcast over the 54-860 Mhz spectrum. These signals often traverse several miles of the coaxial plant while crossing multiple amplification and splitting stages before reaching the customer endpoint. To maintain acceptable signal-to-noise ratio up to the very last subscriber on the cable plant, bandwidth beyond 860 Mhz cannot be effectively utilized in the current design.

The broadcast nature of the downstream plant also raises privacy concerns, especially in commercial applications, because a given customer's data is visible at every endpoint on the cable plant. To mitigate privacy concerns, HFC encryption techniques such as BPI have been added. These techniques introduce additional overhead and complexity.

In the upstream direction, bandwidth between 5 and 42 Mhz is shared across all subscribers. Due to ingress noise issues and interference with other signals, the actual usable portion of the upstream spectrum is further limited to a fraction of this spectrum. One of the methods to access this shared media channel is through a media access protocol defined in the DOCSIS standard. Endpoint cable modems send requests for transmission to a CMTS at the headend which grants permission for accessing unused time-slots in the upstream frame. If a request is

denied, as might happen during periods of heavy load, the cable modem continues to retry until a request is granted. Such a centrally controlled reservation based access scheme leads to large and unpredictable delays making SLA commitments very difficult to meet. While DOCSIS 1.1 allows for prioritized access to time-sensitive traffic such as voice, the delays could still be on the order of tens of milliseconds.

The nature of the initial DOCSIS technology (1.0 and 1.1) is that symmetrical services are nearly impossible and are not possible at all as demand increases. Data flow asymmetry makes it very difficult for DOCSIS-based access systems to penetrate the commercial market.

In summary, the current cable plant architecture has bandwidth and performance limitations that make it difficult for use in commercial applications resulting in low cable modem penetration.

NEXT-GENERATION ARCHITECTURE

The architecture of broadcast cable networks is very similar to that of Ethernet local area networks (LANs) that have been successfully deployed in the enterprise environment for many years. Ethernet LANs have evolved over time to meet the growing needs of the enterprise in terms of increasing bandwidth and applications requirements. The next-generation cable network architecture proposed here for cable networks draws heavily from the lessons learned from the evolution of Ethernet.

Evolution of Enterprise Ethernet LANs

First-generation Ethernet LANs are based on a shared bus architecture where all stations contend for access to the bus using the CSMA/CD access protocol. As the number of

stations and traffic increased, excessive packet collisions cause significant retransmissions that lead to severe delays and throughput degradation.

To alleviate this situation, second-generation architectures segment the shared bus using Ethernet hubs to reduce the serving group size on each segment. However, as bandwidth intensive, multi-service applications continued to emerge, even this architecture can be inadequate. While the hub based architectures allowed some degree of control on managing the overall utilization of a shared Ethernet segment, enforcing Quality-of-Service (QoS) for multi-service applications such as voice, data, and video is extremely difficult as packet collisions cannot be completely eliminated.

The only way to enforce true QoS is to eliminate the shared nature of the network and transition to a point-to-point switched network where every end-user gets a dedicated network link. This recognition leads to the creation of third-generation Ethernet technology - Switched Ethernet - that uses a store-and-forward switching paradigm for forwarding packets.

Switched Ethernet enables several capabilities besides allowing dedicated bandwidth per user. Switching allows packets to be directed only to those endpoints to which they are addressed rather than broadcast to everyone on the network. This eliminates privacy concerns of shared media networks. Through the 802.1q virtual LAN (VLAN) header extension, virtual private network services can be enabled at Layer 2 and the network can be logically segregated into multiple VLAN domains. Enterprises are using this capability effectively to create VLANs for each administrative domain within the corporation on a common switched Ethernet network. The VLAN header also has a priority field that can be used to classify

packets and to service them at different priority levels. Modern Ethernet switches offer priority queuing at each port so that end-to-end delays for critical traffic can be managed more effectively. With adequate buffering at each switching point, the store-and-forward paradigm also allows the network to absorb traffic bursts much more efficiently without losing packets and degrading throughput. Thus switching enables QoS and multi-service capabilities in the network.

Ethernet Switching over Cable Networks

Current cable networks are analogous to first-generation Ethernet LANs – both use a shared bus architecture. Node splitting and the mini fiber-node concept for segmenting the cable network to support smaller serving group sizes is analogous to second-generation hub-based Ethernet LANs.

Due to the inherent deficiencies of the shared nature of existing HFC architectures, the existing network designs cannot scale to meet the needs of the commercial market segment. In order to enable QoS, true multi-service capabilities, and to ensure privacy for commercial customers, a new architecture is needed. Similar to the latest Ethernet technology, the new architecture must have a switching element that is capable of providing secure data to the end user. This new architecture must also deliver all of the services that the commercial customers demand. Clearly, unless the cable operator is willing to give up existing spectrum for the commercial customer, new bandwidth must be activated on the HFC network. The next natural evolution in HFC is to go beyond the existing 860 Mhz spectrum limit. Since continuing to grow a shared network is not feasible, a new data channel along with a switched network overlay must be created on the cable plant. The coaxial cable medium has usable spectrum beyond the 860 Mhz spectrum used by existing systems. The

proposed design makes this previously unused spectrum usable by breaking the tree-and-branch coaxial network into a number of point-to-point coaxial segments interlinked by switching elements. Since each point-to-point coaxial segment is only a few hundred feet long, additional bandwidth beyond 860 Mhz can be effectively utilized for transmitting and receiving high data rate signals while maintaining acceptable signal characteristics. Each coaxial segment can be viewed as a point-to-point link over which the new data channel operates. Data is recovered and relayed at each endpoint. Synchronization is required only link-to-link, eliminating the need to synchronize with a central headend control device and the corresponding end-to-end ranging operations typical of DOCSIS systems. The existing services below 860 Mhz coexist on the same physical cable plant and continue to operate as before.

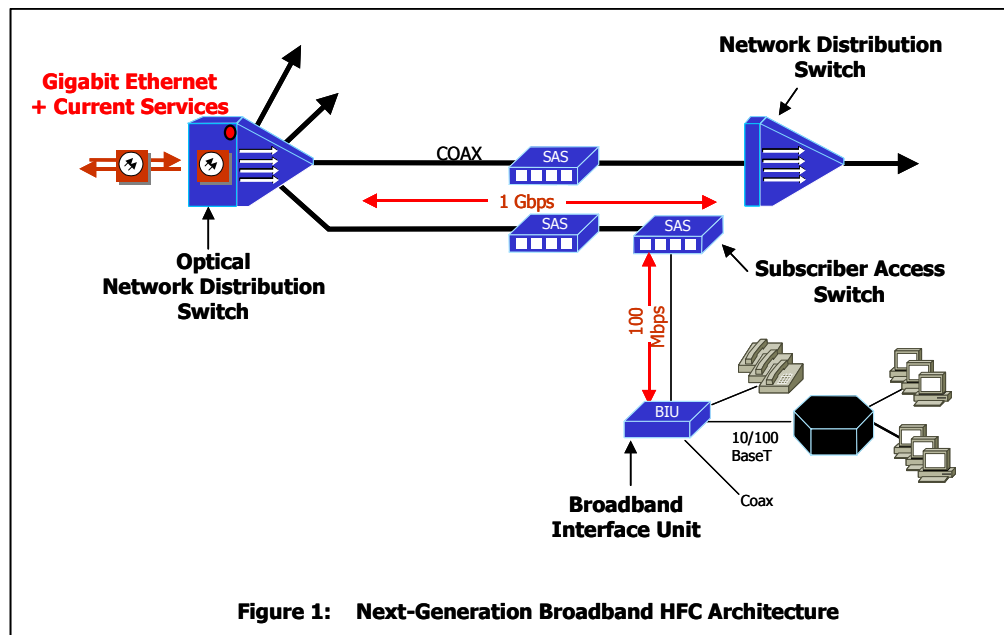
The best choice for the switching technology to be used in this new HFC architecture is Ethernet. Using Ethernet switching in the access network allows for a seamless interconnection of enterprise LANs to metro-optical networks, both of which use Ethernet switching. Additionally, with the use of Ethernet, the headend interconnection architecture is greatly simplified. Since routers support standard Ethernet interfaces, the fiber facilities carrying Ethernet signals from the HFC optical node can be directly terminated on headend routers without requiring any specialized equipment. Although specifically tailored for HFC, the Ethernet segment between the headend and the customer premises functions like a standard Ethernet LAN segment to the headend access router. Since IP over Ethernet services and applications are ubiquitously deployed, this architecture allows the cable companies to readily deliver a rich suite of IP services to their business customers using proven techniques.

Figure 1 shows the next-generation broadband architecture that uses Ethernet switching over the HFC network.

The Broadband Interface Unit (BIU) is the customer premises device that functions as an intelligent IAD. It supports 10/100BaseT interfaces to connect Ethernet attached switches or routers, a coaxial interface for TV and cable modem services operating below the 860 Mhz spectrum and a T1 or POTS interfaces for telephony applications.

available from Frame Relay and ATM networks. Outgoing Ethernet packets are modulated and transmitted over a 100 Mbps upstream channel on the coaxial drop. Ethernet packets are received at the BIU from the coaxial drop over a separate 100 Mbps downstream channel.

The Subscriber Access Switch (SAS) on the cable network combines the operation of a conventional Tap for existing services below 860 Mhz, and HFC Ethernet switching



Traffic is classified and prioritized according to the services provisioned at the BIU. For example, voice over IP (VoIP) traffic is given the highest priority to minimize queuing delays. Mission-critical VPN traffic can be given the next level of priority while Internet bound traffic can be treated at an even lower priority level. The priority fields are appropriately set on an HFC header as well as in the IP header Type-of-Service (TOS) field so that the priority treatment can be extended through the access, metro, and wide-area networks. Traffic policing and shaping can also be enabled at the BIU to support Committed Information Rate style data services similar to those

functions for the high-speed data channels. It supports 100 Mbps full duplex subscriber data channels from customer premises and 1 Gbps full duplex trunk channels that connect to other upstream and downstream SASs over coaxial trunk cables. The Ethernet switching function supports priority queuing and weighted round-robin queuing so that multi-service traffic can be treated appropriately and all subscribers get a fair delay treatment independent of their distance from the headend. The SAS also supports a flow control scheme to avoid packet loss in periods of traffic congestion. The SAS also supports a bypass capability so that when local failures happen, all signals from the upstream trunk

cable are directly coupled to the downstream trunk cable thereby maintaining service to downstream subscribers.

The Network Distribution Switch (NDS) provides the functions of a distribution amplifier for signals below 860 Mhz as well as trunk-to-trunk switching function between the 1 Gbps trunk channels. The optical NDS replaces a conventional optical node and combines the signals below 860 Mhz from the headend with the Gigabit Ethernet signals from the headend router. Switching functions, QoS, and flow control capabilities available at the NDS and the optical NDS are similar to those at the SAS.

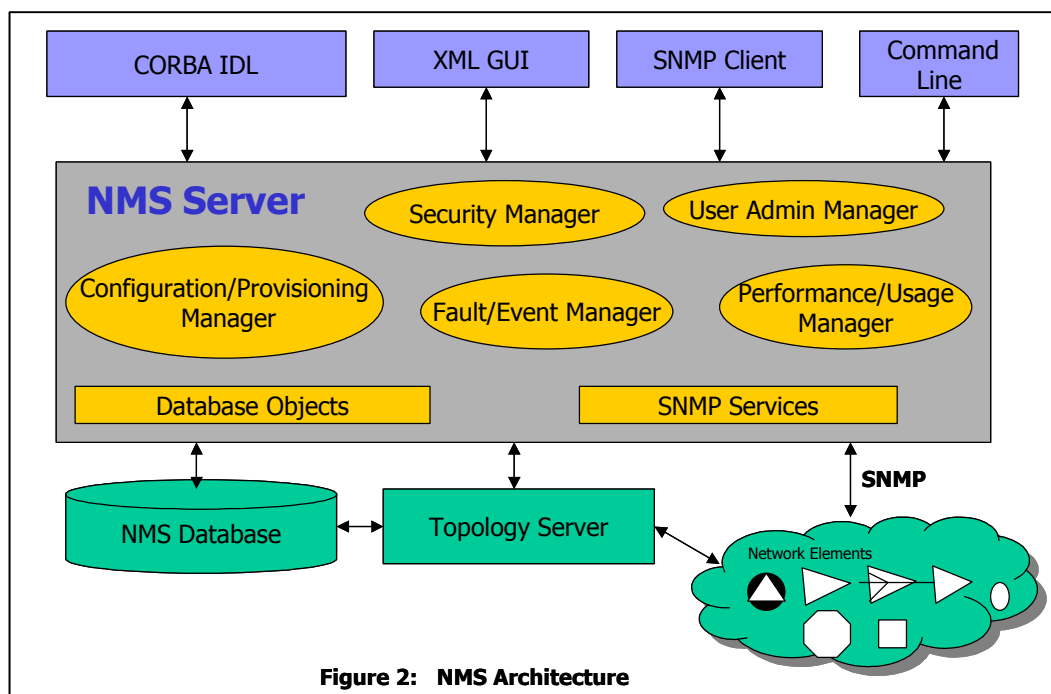
By augmenting the existing HFC infrastructure with 100 Mbps and 1 Gbps Ethernet switching functions, and multi-service QoS capabilities, the proposed architecture provides all the benefits of an all fiber solution at a significantly lower cost. By adopting this architecture cable companies can leverage the ubiquitous nature of the HFC plant effectively.

NETWORK MANAGEMENT

Network management plays a key role in maintaining smooth operation of the network and in ensuring that the network is able to achieve the high standards of reliability demanded by commercial customers. The goal is to be able to remotely manage the network from a local or regional Network Operations Center (NOC) and to fully automate the functions of fault, configuration, and performance management so that truck rolls and other manual operations can be eliminated to the extent possible.

Each element in the new switched Ethernet HFC network is an IP addressable entity supporting an SNMP management agent and protocol. The SNMP management information base (MIB) supports all the operations necessary to manage the network element.

Figure 2 shows the architecture of the Network Management System (NMS). It includes an NMS applications server that hosts management applications and services, a topology server that discovers and maintains



the network topology, and a database server that serves as the repository of network management information. The system supports a number of client interfaces including XML, SNMP, command line, and a Corba IDL for ease of OSS integration. This architecture provides maximum flexibility in distributing and scaling the management functions across multiple hardware platforms and work centers. The goal of the NMS is to integrate as many traditional back office processes as possible. It is only with full integration that the new HFC network will be able to serve the commercial market with timely and reliable services with availabilities at par with alternative fiber based delivery systems. The key management applications supported by the NMS platform are described below.

Configuration Management

When a new element is inserted into the network, it announces its presence and is acknowledged. The appropriate configuration files are then downloaded to the new element. The path taken by the announcement is captured by the topology server which allows it to discover the exact location of the network element. The information about a new element discovery and its location on the network topology is made available to various management applications. At this point, the network element becomes an actively managed element and is subject to all management operations. Thus the initial discovery and configuration of network elements is a fully automated process requiring no manual steps other than physical installation.

Fault Management

The network elements support a broad range of SNMP traps that communicate status changes of various types. The network administrator has the ability to turn these traps

on or off and to direct them to specific operators. Alarms can be classified as minor, warning, or critical and can be used to generate and assign trouble tickets to various technicians. The fault management application also allows the network administrators to run a full suite of diagnostics and to perform various trouble isolation tests remotely to quickly diagnose fault conditions. The network elements themselves support a built-in self-test (BIST) capability for carrying out successful initialization tests. The network administrator can initiate BIST capability for fault diagnosis when required. The network administrator has the ability to remotely reset devices to clear fault conditions. These operations ensure that truck rolls are minimized and the mean time to detect and repair faults is as short as possible thereby improving network availability.

Performance Management

The NMS is capable of collecting a number of statistics related to the physical layer and data link layer. At the physical layer the operator can remotely collect and display power levels of the RF signals. In this system, gain and slope characteristics of the amplifiers can be adjusted either automatically or under operator control based on RF measurements. At the data link level, statistics on frame and byte counts and dropped packets can be collected by service classes. In addition to validating SLA performance, the various statistics can be used for accounting purposes and fed to the billing systems after appropriate processing. The network administrator has the ability to set what statistics are to be collected and at what intervals. As with other state of the art management platforms, the NMS system described is able to filter and present custom views of the network and its performance to individual users responsible for specific deliverables.

ENABLING COMMERCIAL SERVICES

Enabling services for SMB customers on the proposed broadband network infrastructure is key to new revenue generation. The BIU platform becomes an intelligent service edge capable of delivering value-added functionality to the end user. Some of these services are described below.

IP Telephony

Telephony service is part of the core set of services demanded by commercial customers. While some cable companies have been offering residential telephony services using circuit-switched technology, this technology is expensive and difficult to scale due to its inefficient use of the frequency spectrum. SMB customers require upwards of four telephone lines on average and therefore the aggregate number of lines could easily become large with sufficient penetration rates. Cable modem technology using DOCSIS 1.1 standards would allow more telephone lines to be supported on the cable network by packetizing and compressing voice traffic. However, while adequate for residential customers, voice compression degrades performance and may not be acceptable to business customers.

The QoS enabled switched Ethernet architecture described in this paper presents an ideal platform for supporting IP telephony. The 100 Mbps / 1 Gbps link bandwidth is adequate to support a large number of telephone lines without requiring voice compression. Furthermore, the priority treatment at each hop ensures that end-to-end delays experienced by voice packets over the cable access network are in the sub-millisecond range from the premises to headend. Therefore, next-generation broadband technology has the potential to accelerate the deployment of toll-quality IP telephony.

From a telephony perspective, the intelligent BIU may function as an IAD by providing T1 or POTS interfaces to support on site PBX or analog phones. Voice traffic is packetized into VoIP format and transmitted over the Ethernet interface. Signaling information is transmitted using standard signaling protocols such as H.323, MGCP or SIP. On the network side, the BIU communicates with a VoIP gateway and a softswitch which shunts traffic to and from the PSTN. Value-added services such as Centrex can be offered by a feature server connected to the softswitch. Recent advances in the softswitch architecture are making it possible to support key business services without requiring an expensive Class 5 switch.

Internet Access

Many SMBs require a high speed Internet access service. This service can also be readily enabled by the intelligent BIU. Most cable companies have shown a preference for using policy based routing (PBR) as the mechanism to support multiple ISP access. The BIU can be provisioned with the appropriate routing and QoS policies for different source IP subnet and MAC address combinations based on the levels of service offered by different ISPs.

In addition to PBR, the BIU can support other options for enabling multiple ISP access. For example Layer 2 VLAN technology can also be used for this purpose. Each ISP subnet can be represented by its own VLAN domain and the BIU can be provisioned to map specific source MAC addresses to VLAN domains based on the ISP selected by the subscriber. Different QoS policies can be applied within each VLAN domain if the ISP wishes to offer different grades of service to its subscribers.

The BIU's ability to uniquely classify packet flows using any combination of IP address, TCP/UDP port, and MAC address, can be exploited to support fine-grained QoS features linked to the various applications and services offered by the ISP.

Virtual Private Networking

Small and medium size businesses require VPN services to link branch offices and to support remote access for telecommuters and work-at-home employees. VPN services can be enabled at the BIU by supporting various forms of tunneling mechanisms. Site-to-site VPN connectivity can be provided either by Layer 2 VLANs or by Layer 3 tunneling mechanisms such as IPSec. BGP/MPLS VPNs can also be enabled at the BIU. For telecommuter access, L2TP tunneling can be supported from the BIU without requiring any specialized client software. All forms of tunnels can be provisioned with QoS attributes to support multiple grades of service.

Storage Area Networking

The high bandwidth and low latency characteristics of a Gigabit Ethernet infrastructure can be exploited by the cable companies to offer new services such as storage area networking (SAN) to small and medium businesses. The incumbent service providers cannot offer such a service today because they lack a true broadband infrastructure.

SAN technology has emerged as a mechanism to provide servers with shared access to storage resources over a high-speed network. Currently businesses spend a lot of money to build dedicated SANs using highly specialized fiber channel networking protocols. By outsourcing this function, businesses can save significant amounts of money in both SAN hardware and software as

well as in the IT staff that would manage the SAN infrastructure.

With the emergence of Gigabit Ethernet, there is a move to consolidate SAN applications over an IP network so that specialized networks can be avoided. The recent progress in SCSI-over-IP standards, and the development of iSCSI host adapters for servers and routers promises to bring SAN and IP networking together. With a ubiquitous broadband network, cable companies can exploit this opportunity and position themselves as the SAN service provider of choice to small and medium businesses in major metropolitan areas.

SERVICE DELIVERY PLATFORM

With the creation of a high-speed, symmetrical broadband access network, broadband services can be delivered from a metro service delivery center that attaches to the cable company's metro-optical network. This architecture is shown in Figure 3.

The metro service delivery center hosts a number of servers and gateways that support the kinds of services discussed earlier. For IP telephony, media gateways, softswitches, and feature servers provide value-added services and connectivity to the PSTN. Gateway routers provide connections to multiple ISPs. Content and storage servers are used to provide hosting and storage services.

In addition to servers and gateways, the metro service delivery center also hosts a service delivery platform (SDP) that supports provisioning, billing, customer care, and other back-office functions necessary to ensure smooth delivery of services. The SDP interfaces with a broadband services portal that is used by customers to order and provision services online and provides the interfaces to provision and control the

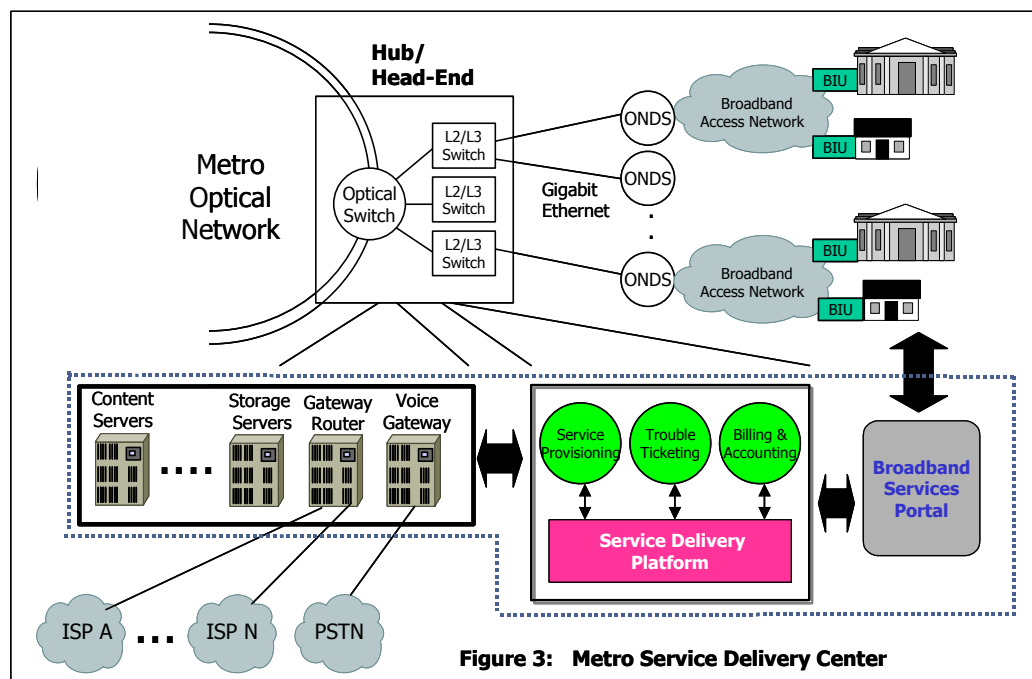
gateways and servers that provide service functionality.

The main challenges in delivering services to end-users are:

- Ensuring that services can be ordered and provisioned in near real-time,
- Ensuring that new services can be rapidly created and deployed without requiring major software upgrades on network elements,
- Providing customers with a level of control in ordering and self-provisioning services and in monitoring their service status.

and the workflow, thereby allowing existing databases to be used and preventing unnecessary duplication of databases. By adopting a Java based framework similar to the NMS, the DEN applications and databases can be fully distributed leading to operational flexibility and scalability. Figure 4 shows the proposed DEN-based SDP architecture.

The main platform components of the SDP architecture are described below.



A Directory Enabled Networking (DEN) architecture for the SDP can properly address these challenges. DEN has been used successfully in enterprise networks to manage networks, service agents, and the client interaction with services in a policy-driven framework. By properly associating subscribers with network elements and services, the DEN model allows services to be rapidly created and provisioned. DEN decouples service data from the service logic

Service Creation Platform

The service creation platform (SCP) allows service providers to define the service logic associated with offering new services. The definition of service logic includes the description of a service template and a description of the provisioning workflow associated with the service. A service template describes all the characteristics of the service such as the service elements used by the

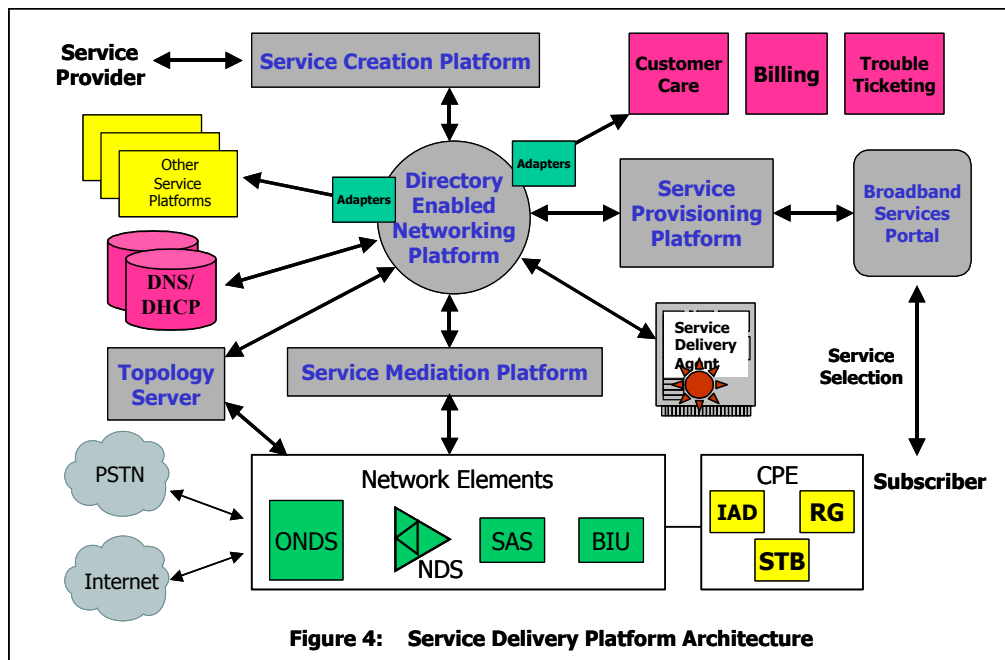
service and a list of service parameters that need to be configured on each service element. The provisioning workflow describes all the tasks and the sequence in which they need to be executed to fully enable the service for the subscriber. Service creation can be a complex task requiring highly skilled personnel. A well-developed SCP is able to hide this complexity by providing the cable operator with a GUI and service creation object classes that makes service creation a simple and easy exercise.

LDAP Directory

The directory is simply a description of an information model and a data model that captures

based LDAP interface. Applications can retrieve any information related to subscribers, their services and their attached network elements through a unified LDAP query to the directory without knowing where the data is physically stored. This capability eliminates significant complexity in application design.

A key application of the directory is to interface to other existing back-office systems used for customer care and billing. This is done through the use of appropriate “adapter” technologies that allow data to flow between the systems. Through these adapters, existing customer databases can be accessed for customer account information, and the appropriate billing data can be passed to existing billing systems.



the relationship between subscribers, the network elements they are attached to, and the services to which they subscribe. The service templates and the provisioning workflows are also captured in the directory. The data model captures the data associated with a subscriber’s service. The actual data can reside anywhere on the network in physical databases. The information model is exposed to external applications through a standards

Service Provisioning Platform

The service provisioning platform (SPP) is the major workhorse for executing all flow-through provisioning tasks. It includes a powerful workflow engine that can execute the tasks defined in the provisioning workflow, undo tasks when some of the provisioning steps fail, and update the appropriate back-office systems that support other service management functions when the

provisioning steps are successfully executed. The SPP is a high performance platform capable of executing several provisioning requests simultaneously.

The SPP gathers subscriber provisioning data from the broadband services portal and builds the association with network and service elements by accessing the directory information model. It parses out the provisioning data according to the impacted service elements and then proceeds to carry out the provisioning tasks.

Service Mediation Platform

The service mediation platform (SMP) is used to manage the network elements during service provisioning. It hosts the service MIB that captures the management semantics of the particular service being provisioned. It takes provisioning commands from the SPP and translates them into the appropriate SNMP set requests to provision the device MIB on the BIU. By separating the service MIB from the device MIB, this architecture allows multiple services to coexist on the device without hard-wiring service level semantics into the device. By supporting a Java Virtual Machine (JVM) based service execution environment on the BIU, service logic can be downloaded at provisioning time and executed on the fly. This architecture embodies a truly dynamic service delivery environment.

A major function of the SMP is to perform admission control on both real-time services such as IP telephony and provisioned services such as VPNs. This ensures that adequate capacity is available in the network to support the service at the requested QoS levels so that customer SLA commitments can be honored.

SUMMARY

The next-generation broadband HFC architecture described in this paper unlocks the full potential of the most widely deployed asset of a cable company – the HFC network. By adopting a point-to-point switched Ethernet architecture using spectrum beyond 860 Mhz, previously unused bandwidth can be effectively utilized to create a high capacity two-way packet data channel. This multi-service QoS enabled channel can be used to deliver a wide range of IP based broadband services that were not possible before. In this new broadband network, services can be created on the fly and made available to the subscribers in near real-time through a sophisticated service delivery platform. Armed with such a state-of-the-art broadband infrastructure, cable companies can position themselves strongly in the marketplace, giving end-users affordable access to true broadband services, for the first time.

HUB, BRIDGE, SWITCH, ROUTER: DEFINING THE FUNCTIONS OF THESE DIFFERENT DEVICES

Doug Jones
YAS Broadband Ventures

Abstract

Hubs, Bridges, Switches and Routers all play roles in creating LANs and IP networks. This paper discusses the networking characteristics these different devices.

Thirty years ago it was antenna design. Twenty years ago it was amplifier design. Ten years ago it was fiber optics. Now, it's the Internet. The basic building blocks of intranets and the Internet include the hub, bridge, switch, and router. These devices all have specific functions and roles to play when piecing together an IP network. The intent of this paper is to introduce each piece of equipment, describe what it does, and how it does it.

The first three devices; hub, bridge, and switch, are very closely related. They all operate on Ethernet frames. The fourth device, a router, operates on IP packets. The difference between a frame and a packet is semantic, having to do with which layer in the OSI protocol model a particular protocol operates. Regardless, both a frame and a packet are standardized sets of bits used to represent data.

This paper will begin with a brief description of the basics of Ethernet and how it works. This introduction is necessary to draw the distinctions between a hub, bridge, and switch. Next, the paper will briefly describe the Internetworking Protocol (IP) and how a router works.

Ethernet Background

Ethernet was invented by Bob Metcalfe in 1972 at the Xerox PARC (Palo Alto

Research Center). This first Ethernet operated at about 3 Megabits per second (Mbps). In 1980, the first 10 Mbps Ethernet standard was published by DEC, Intel, and Xerox and is known as DIX Ethernet. At this point, the IEEE 802 committee, responsible for LAN/MAN standards development, took up the original DIX specification and used it as the basis for an IEEE standard. This work was completed in 1985 by the IEEE 802.3 subcommittee and is titled the *IEEE 802.3 Carrier Sense Multiple Access with Collision Detection (CSMA/CD) Access Method and Physical Layer Specifications*. Note that the word Ethernet does not appear in this title, however, both the IEEE 802.3 protocol and DIX are simply referred to as Ethernet.

The IEEE 802.3 standard is the official Ethernet standard, and it differs slightly from the DIX Ethernet standard, though the two are backward compatible. An Ethernet frame is shown in Figure 1. The difference lies in the Type/Length field. In DIX Ethernet, this field is used to identify the type of protocol carried in the payload. In IEEE 802.3, this field indicates the length of the payload. Though this one field is used for two purposes, it is possible to determine how a particular frame is encoded. When originally developed, Xerox made sure that none of the protocol identifiers had a value of less than 1500; hence, if the value is greater than 1500 it is a DIX frame. If the value is 1500 or less, then it is an IEEE 802.3 frame. In the case of an IEEE 802.3 frame, there are additional fields (LLC fields) in the payload that identify the type of protocol being carried in the payload.



Figure 1

The source and destination addresses are both 48 bits in length, and are commonly referred to as “MAC addresses.” The term MAC stands for Media Access Control, which is a generic term for a data link layer protocol, which Ethernet is. Note the payload is up to 1500 bytes. While an Ethernet frame can carry many types of payload, later on this paper will address the case when that payload is an IP packet.

The Ethernet Collision Domain

Ethernet networks originally operated over coaxial cable and multiple devices were connected to the same cable segment. Because multiple devices connected to this cable, there was the possibility of two devices transmitting at the same time. When two (or more) devices transmitted at the same time, those transmissions would collide and the data would be lost.

The IEEE 802.3 name for Ethernet, CSMA/CD, gives a hint at how Ethernet attempts to minimize collisions. Before transmitting on the cable, a device first “listens” on that cable (carrier sense) for existing transmissions. If the device senses the cable is busy, it does not transmit. If the device senses the cable is idle, it will transmit. The worse case scenario is when devices on opposite ends of a cable segment sense the cable is idle and begin transmitting at the same time. Due to propagation delay, this collision takes the longest time to detect.

Ethernet was designed to guarantee that a transmitting station could tell whether its transmission had failed due to a collision. In order to meet this guarantee, the length of cable had to be limited and the minimum

frame size had to be specified. Given that the maximum cable length was 2.5 kilometers, given the speed of electricity on the cable, and given a transmission speed of 10 Mbps, it might take as long as 512 bit times to detect a collision, hence the minimum frame size of 64 bytes (64 bytes x 8 bits/byte = 512 bits).

As the number of devices on a cable segment increases, so does the traffic and the probability of a collision occurring. Collisions mean that data needs to be retransmitted and hence is an inefficient usage of the available bandwidth. One way to keep the number of collisions low was to keep the length of cable short.

What’s an Ethernet Hub

An Ethernet hub is best described as an Ethernet repeater. That is, when a bit is received on a hub port, that bit is repeated onto every other port on the hub. On the one hand, this is a dumb, low cost device. On the other hand, a hub increases the size of the collision domain.

Each cable plugging into a hub is considered a “LAN segment.” Each LAN segment is a collision domain. A hub, since it just repeats bits onto every segment, increases the size of the collision domain seen by all packets.

Consider the case of a 4-port hub. (There is nothing magical about 4 ports, this is just a standard configuration for commercially available hubs.) There are LAN segments connected to three of the 4 ports, call them segments 1, 2, and 3. A device connected to LAN segment 1 has data to transmit. This

device follows the CSMA/CD MAC protocol by “listening” on its cable segment and, determining it is idle, begins transmitting. The hub receives these bits and relays them onto LAN segments 2 and 3 without bothering to determine if those segments are idle. Hence, if other devices are transmitting on those segments, there are collisions.

Hubs are low cost devices intended for the quick interconnection of low traffic LAN segments. When used to connect high traffic LAN segments, hubs increase the probability of collisions, lost data, and retransmission on all of those segments. Cascading hubs should be discouraged as this just compounds the problem.

What’s an Ethernet Bridge

A bridge is a lot like a hub, except this device respects the collision domain of the individual LAN segments connected to it.

Consider again the 4 port device, but this time a bridge. When a bridge receives a transmission, it buffers (stores in memory) the entire frame. Here is a first key differentiation from a hub. A bridge can determine what bits make up a frame, as opposed to simply repeating bits. Before the frame is sent onto other segments, the bridge “listens” on each of those segments for idle. A second key differentiator is that a bridge implements the CSMA/CD protocol on each of its ports, while a hub does not.

Therefore a bridge is a higher cost device. It has to contain both extra memory and the logic to perform the MAC protocol. However, a bridge does not increase the size of a collision domain like a hub does.

What’s an Ethernet Switch

A switch is a lot like a bridge, in fact, another name for a switch is a “learning” bridge. As frames pass through the various ports of a switch, the switch learns which

devices are connected to the LAN segments and stores their Ethernet MAC addresses in memory.

To learn how the switch operates, continue with the above example but place two devices on each LAN segment. As Figure 2 shows, segment 1 has devices A and B, segment 2 has devices C and D, and segment 3 has devices E and F.

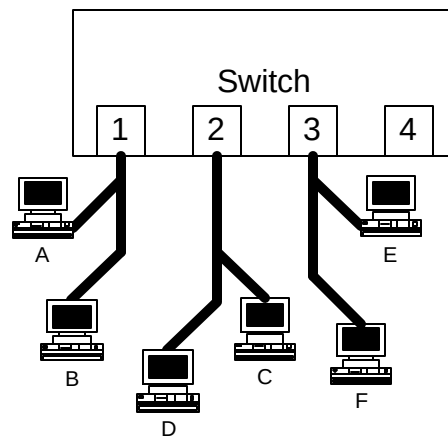


Figure 2

In this case, if device A sends a frame destined for device B, even though that frame is received at the switch, the switch will not forward the frame onto the other LAN segments because the switch knows that A and B are on the same LAN segment.

Likewise, if E sends a frame to B, the switch will only transmit that frame on segment 1 and will not transmit that frame on segment 2. Like a bridge, a switch also implements the CSMA/CD protocol.

If the switch receives a frame and does not recognize the destination MAC address (i.e., broadcast frames, the first time a device transmits, frames destined to other places, etc.), that frame will be forwarded onto all other LAN segments.

More sophisticated switches allow the creation of Virtual LANs (VLANs) where a subset of ports can be grouped together into a

frame format will differ, the IP packet carried in those frames will remain the same.

On each LAN segment, the router port

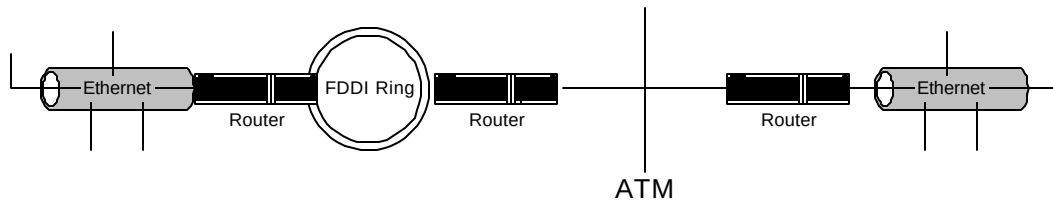


Figure 3

logical “switch. For instance, a common commercial-grade switch configuration is to have 96 ports. A subset of these ports can be logically grouped together to create smaller switch domains.

IP Background

Up to now, this paper has discussed Ethernet frames and the role of hubs, bridges, and switches in maintaining the most efficient collision domain on LAN segments. IP routers are an entirely different type of device. True, routers have Ethernet interfaces (10Base-T, 100Base-T, etc.), but IP has a different function than Ethernet.

Ethernet is a point-to-point protocol, and IP is an end-to-end protocol. While both use addresses, an Ethernet address only has meaning on a LAN segment whereas an IP address has global meaning.

As shown in Figure 3, routers can be connected with many different types of LANs. Ethernet, FDDI, and ATM are shown but there are many others. On each of these LANs, specific frame formats are used for carrying data. Just as the Ethernet frame in Figure 1 has a specific format, FDDI and ATM frames have their own formats and rules. All of these frames can carry IP packets. While on each LAN segment the

will have a MAC address that is specific for the individual LAN type. As the IP packet goes from router to router (and traverses the LANs), the frames carrying that packet will change depending on the technology of the underlying LAN; however, the IP packet carried in the frame will not be changed.

What’s an IP Router

A router has two main functions, IP routing and packet forwarding.

IP routing is the process whereby routers exchange route information so they know where to forward IP packets. Route information is exchanged by using a variety of protocols, including:

- RIP (Routing Information Protocol)
- OSPF (Open Shortest Path First)
- BGP (Border Gateway Protocol)
- Many others

Routing protocols are a study in and of themselves. Suffice it to say routers are continually exchanging route information because routes across the Internet are continually changing for a number of factors. For example, different paths through the Internet can be congested or out-of-service at any given time. Routing protocols identify these bottlenecks and provide routers with “up to the second” information to keep the traffic moving.

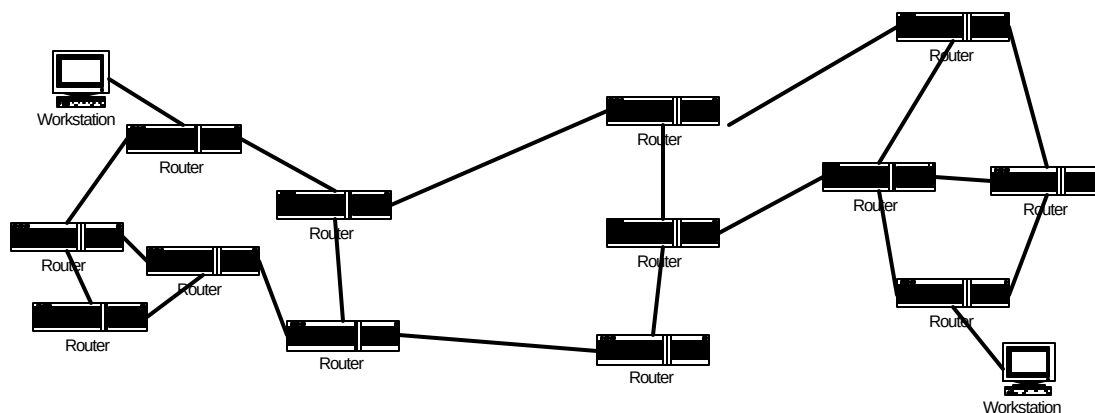


Figure 4

Consider the case of a single large email message traversing the Internet. This message is too big to fit into one IP packet and therefore must be fragmented into many IP packets. In going from one email account to another, the individual packets of this email message may take different paths across the Internet. As these packets are received at the other end, even out of order, higher layer protocols are responsible for reassembling them into the original email (in this specific case, that protocol is TCP, or the Transmission Control Protocol).

These self-healing and redundant properties make the Internet very robust. The original Internet was created under the guidance of the U.S. Military in the 1960s and one of the design criteria was to remain operational even in the event of massive destruction of parts of the network (i.e., nuclear attack).

The routing protocols exchange route tables, essentially lists of information that identify paths through the Internet to particular destinations. When an IP packet arrives at a router, the IP address in that packet is used to index into the route table and the route information is used to decide how to forward that packet. While Figure 3 was an example of a simple connection between routers, a more realistic network design is given in Figure 4.

In Figure 4, there are over a dozen possible paths to get data between the two workstations. The routing protocol calculates the most efficient path between the two workstations and the routers in between forward packets along that path. These decisions are made on a hop-by-hop basis. Based on its route table, each router along the path makes an autonomous decision about the best path to forward the packet. In the case of one router (or LAN segment) going out of service, the routing protocol quickly (on the order of seconds) recalculates the new most efficient path for packets to take and propagates this information among the affected routers. While some packets may be dropped during the reconfiguration, the network is self-healing and communications continue.

The difference between a MAC address and an IP address

As written earlier, Ethernet is a point-to-point protocol and IP is an end-to-end protocol. This means that as a frame traverses LAN segments, its source and destination addresses will be changed, however, the IP addresses in the IP packet carried by that frame will remain the same.

frame, and sets the source MAC address equal to that of the workstation and the destination MAC address to that of the first router. The workstation then puts this frame on the Ethernet link according to the CSMA/CD protocol rules. The router “sees” the frame on the LAN segment with destination MAC address set to its own, hence this frame is destined for that router.

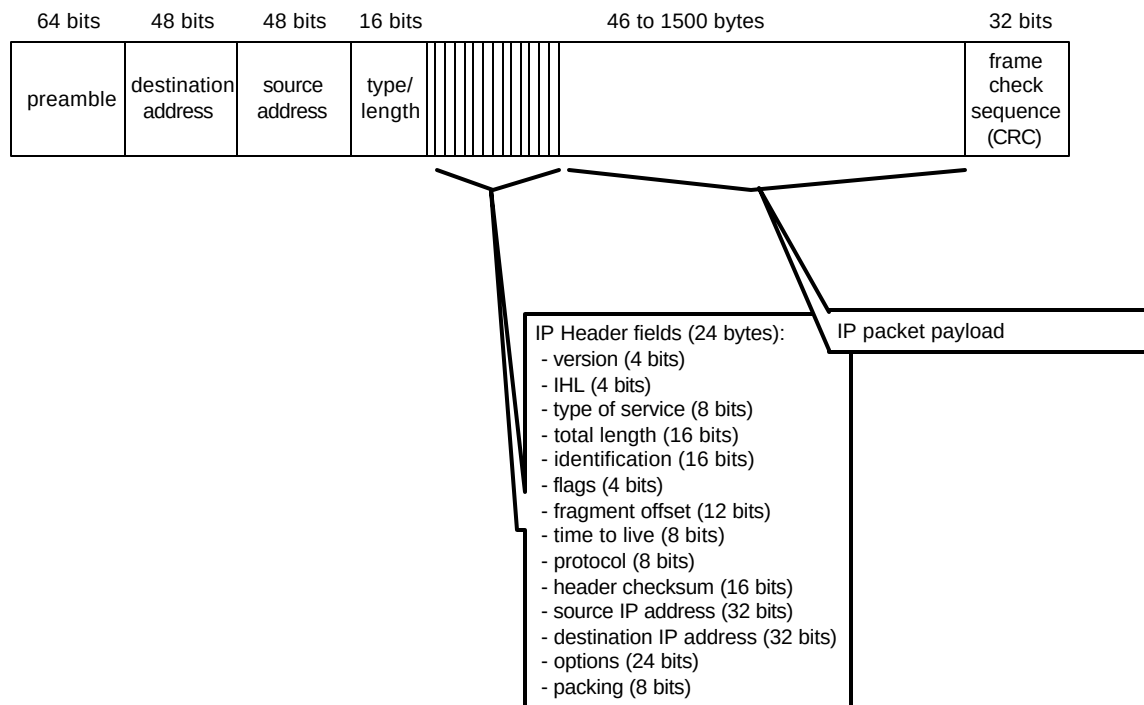


Figure 5

Figure 5 shows a redrawing of Figure 1, except the payload of the Ethernet frame contains an IP packet. Consider the case when this Ethernet frame (carrying the IP packet) is forwarded from one workstation to another (as diagramed in Figure 4).

The scenario is as follows. An application running on the workstation creates the IP packet and sends it to its Ethernet LAN interface for transmission. The IP packet has the source IP address set to this workstation and the destination IP address set to the other workstation. The Ethernet LAN interface on the workstation encapsulates the IP packet in an Ethernet

Although other devices on the LAN segment see the frame, they would discard this frame because the destination MAC address is for the router. (Note, while Ethernet MAC addresses only have significance on a particular Ethernet segment, each Ethernet MAC address is unique in that it should never be duplicated in another device.)

Now that the router has the Ethernet frame, that router strips off the Ethernet bits and uses just the IP packet (mainly the destination IP address) to make the next-hop packet forwarding decision. When that packet exits the router, the router re-

encapsulates the IP packet in a new Ethernet frame (or whatever protocol is appropriate for that LAN segment, as proposed in Figure 3) with a source MAC address of the router port and the destination MAC address of the next hop router. Note that the IP addresses in the IP packet have not been changed. In this manner, the IP packet makes its way, router-to-router, LAN segment-to-LAN segment, between the two workstations. When finally delivered to the destination workstation, the Ethernet frame will contain the MAC address of the last-hop router in the source address field and the MAC address of the workstation in the destination address field.

Summary

Ethernet and IP are separate protocols that do different things. However, IP packets are routinely encapsulated in Ethernet frames for transmission on LAN segments.

Hubs, bridges, and switches operate in the Ethernet domain and Routers operate in the IP domain. Hubs, bridges, and switches connect Ethernet segments. Hubs are low cost, but increase the collision domain a particular frame is exposed to. Switches have additional cost for memory and logic, but they create more efficient LAN segments.

Routers perform routing and forwarding. Routing involves exchanging route information with neighboring routers on the network. The route information is used to forward individual IP packets. Routing is very dynamic and is a main reason why the Internet is such a robust network.

References:

1. Spurgeon, Charles E.; Ethernet - The Definitive Guide, O'Reilly & Associates, 2000.
2. IEEE 802.3 Carrier Sense Multiple Access with Collision Detection (CSMA/CD) Access Method and Physical Layer Specifications, IEEE 802 Committee, 1998.
3. Perlman, Radia; Interconnections - Bridges and Routers, Addison Wesley, 1992.
4. Huitema, Christian; Routing in the Internet, Prentice-Hall, Inc., 1995.

Author Contact:

Doug Jones, Chief Architect
YAS Broadband Ventures
voice: (303) 661-3823
doug@yas.com

IMPROVING CABLE SYSTEM BANDWIDTH MANAGEMENT BY UTILIZING SUSCRIBER-BASED STORAGE RESOURCES

By

Tim Elliott – Senior Director Systems Architecture for Keen Personal Media, Inc.

Peter Schwartz – Director of Product Management for Keen Personal Media, Inc.

With

Joseph Weber – Technical Consultant

Keen Personal Media, Inc.

Abstract

Cable head-end systems supporting on-demand content delivery applications and other key services operate 24 hours a day, seven days a week. Service networks need the ability to scale performance to handle large volumes of subscriber requests during high demand periods without creating unwanted delays or incurring prohibitive costs. Additionally, cable system operators need solutions that are scalable while maintaining efficient bandwidth utilization as workloads increase.

For these reasons, the addition of local client storage in the set-top box presents tremendous benefits to the system operator. One of the many key features of local storage will be the ability to offer customers true video-on-demand: a value-added service already familiar to the customer. This paper will explore the benefits of managing the delivery of content traffic to several client devices with local storage. The result is a video-on-demand service with higher availability, easier manageability, and greater scalability than previous solutions.

The Content on Demand Service

The cable industry has the unique capability of delivering high quality on-demand and targeted content to its customers, thereby delivering a premium service superior to other media delivery systems. The cable customer desires and now can receive true video-on-demand (VOD) services;

this includes the ability to start a conditional access program of his or her choice at any time, to pause/rewind/fast-forward the program just like a VCR, and in some cases to view the content multiple times within a given time period. The content on demand concept is already familiar to the consumer because it is similar to the activity of renting programs from videotape rental outlets. Cable-based VOD offers the same service but with an increased level of convenience. The frustration associated with going to a rental outlet only to find that the program you want has been sold out, or is otherwise unavailable ceases to be an issue with client-side enabled VOD. Cable VOD also eliminates the specter of late fees that are incurred when returning media to a rental outlet outside of a given timeframe. Therefore the VOD services provided by cable network operators will offer consumers a familiar premium service where the benefits are immediately obvious to subscribers.

This paper will review some of the technical issues addressed by Keen Personal Media's TV4me™ software and hardware solutions to providing VOD and demonstrate that local storage at the client set-top allows the true VOD experience without using excessive amounts of network bandwidth, implementing major infrastructure changes, or creating excessive deployment costs. Note that we will examine only those technical solutions that provide the consumer with the true VOD experience. Since they do not meet the authors' requirements for

true VOD, solutions like certain pay-per-view (PPV) services involving a carousel of staggered start-times and conditional access programs that require the customer to wait for the next start time are not included in this paper. By limiting our examination to the full VOD experience, which removes the delay in start time and enhances the viewing experience with the ability to fast-forward and rewind, we make the assumption that the consumer would be more willing to pay for the premium services that only the full VOD experience can provide.

Technical Solutions for VOD

Clients without local storage

Early trial systems for providing VOD to consumers included two-way, digital set-top boxes capable of decoding MPEG-2 Transport Streams and a network infrastructure to provide direct connections between media servers and set-top boxes. However, the set-top boxes did not include a local storage device sufficient for storing streaming data for any significant length of time. Examples of these early trial systems include the Time Warner Full Service Network (FSN) and the Time Warner Pegasus Phase 2 deployments described by Michael Adams in his book, *Open Cable Architecture*¹. In these examples, a direct communication link is established between a media server and the consumer's particular set-top box. This link requires at least one dedicated forward digital channel assigned to that customer that is capable of carrying 4 to 6 Mbps of bandwidth without interruption. A secure and persistent back-channel connection was also needed for purchase and control of the stream. In order to provide the consumer the ability to fast-forward and rewind the stream, additional dedicated forward digital channels were sometimes used.

Since dedicated channels must be created in real-time to serve each customer, as they demand content, sophisticated switching and server systems need to be added to the system. For each collection of subscribers, there are several requirements: media servers are required to store

and stream the content; forward channel switching systems are needed to establish the dedicated channel to the customer from the media server; and bandwidth must be allocated from a dedicated range of digital channels. Such infrastructure changes affect the head-end, distribution nodes and the set-top box client components of the system.

As illustrated by the Pegasus case study, on average one 6MHz channel can contain 10 digital content streams using QAM modulation. Therefore six dedicated 6 MHz channels (60 streams) were needed to serve 600 customers (where the peak utilization was assumed to be 10%). For a distribution hub with 20,000 homes passed, 2,000 simultaneous streams would need to be served, or 200 digital channels. Due to the limitations of simultaneous streams and media servers that can only handle a finite number of customers, network equipment must be replicated for every group of customers. The true expense for providing VOD to customers in these deployments comes from its inability to scale to serve all customers. Depending on the number of subscribers, the estimated cost of the network changes is \$150-\$200 per VOD stream in addition to the head-end equipment cost of several million dollars. Note that the upgrade is required for all subscribers on this part of the network, regardless of their subscription to the service. In addition, network upgrades are required to provide the capability to serve customers at peak times like weekend evenings. During off-peak times the network bandwidth is underutilized.

The following table summarizes the use of set-top boxes without local storage to provide VOD to customers:

<i>Set-tops without storage</i>	
Advantages	
•	Reduced cost set-top without storage
•	No security risk from local copy

Disadvantages	
• Major network infrastructure changes required to support VOD • Bi-directional communication required throughout the session • A dedicated digital stream to each customer using the service. • Poor scalability and bandwidth utilization	
Cost Per Stream	
• \$1,500 - \$2,000 cost per VOD stream	

Brute Force Method: Store all content at the set-top

Most of the expense and complication of the early systems can be alleviated by using set-top boxes (clients) with local storage (usually in the form of a hard disk drive). Local storage can remove the requirement that a persistent direct connection has to be created between the appropriate media server and each client. In addition, local storage of content can result in more efficient and faster trick-play (fast forward, rewind, etc.) features by removing latency between the user's requests via the remote and the response. In early systems, remote commands had to travel back to the media server in the network. With local storage, the commands are executed locally and therefore without delay. Note that local storage in the set-top provides additional services to the customer in addition to conditional access VOD services. Customers with local storage in their set-top receive all the benefits of personal video recorder (PVR) technology such as pause of live broadcast and EPG-based program recording. PVR devices currently available through retail outlets have already demonstrated that customers are willing to pay for such services. However, the exact value to customers is still to be determined.

We present here a brute-force method of providing true VOD to customers by using local set-top box storage. This model is used only as a comparison to a more efficient system to be described later. In this model, the conditional access content for the VOD service is stored in its entirety at the client set-top box. This content might have been downloaded the night before

during off-peak bandwidth times. When the customer requests (and pays for) the content, it is unlocked from the local storage and available for a finite time period to the customer.

This brute-force method would require the least amount of network bandwidth. Each conditional access program is broadcast once during off-peak times and stored on the local set-top. However this system has two major drawbacks that make it less practical. The first is the large storage requirements within the set-top. Assuming that each hour of stored content requires 2 GB of storage, a significant portion of the local storage would be dedicated to storing the available content. Given that less than 10% of that content would ever be accessed, this would be a very inefficient use of the local storage. The second issue is security. Local copies of the content can be protected using existing scrambling schemes, but having persistent copies of the media in the set-top will be a concern to the content owner.

The cost of implementing such a system includes the additional cost of the storage within the set-top unit. Current hard disk drive costs are \$2.50 - \$3.50 per GB. And, today the most cost efficient capacity is 40 GB. (We can also expect these numbers to continue to drop and capacities to increase over time.) We assume a cost of \$150 - \$180 for local storage, which is added to the cost of a standard set-top. This cost scales linearly with the number of subscribers to the service. Earlier VOD models scaled with anticipated peak load usage and used a network-storage model versus local storage.

Unlike VOD solutions without local storage, major upgrades to the distribution system are not needed in a local storage solution. There also may be some argument that consumers themselves may be willing to absorb some of the additional cost if the set-top was available at retail. We do not make any such assumptions in this analysis.

The following table summarizes the features of the brute-force method of providing COD using local storage in the set-top client:

<i>Storing all programs on the set-top</i>	
Advantages	
• Simpler network infrastructure requirements • Provides PVR services to each customer • Reduced deployment cost	
Disadvantages	
• Security issues with locally stored content • May use a large percentage of the local storage capacity • Number of on demand titles driven by capacity	
Cost Per Subscriber	
• \$1,500 - \$2,000 cost per VOD stream plus \$150 - \$180 for hard drive in set-top box per home passed	

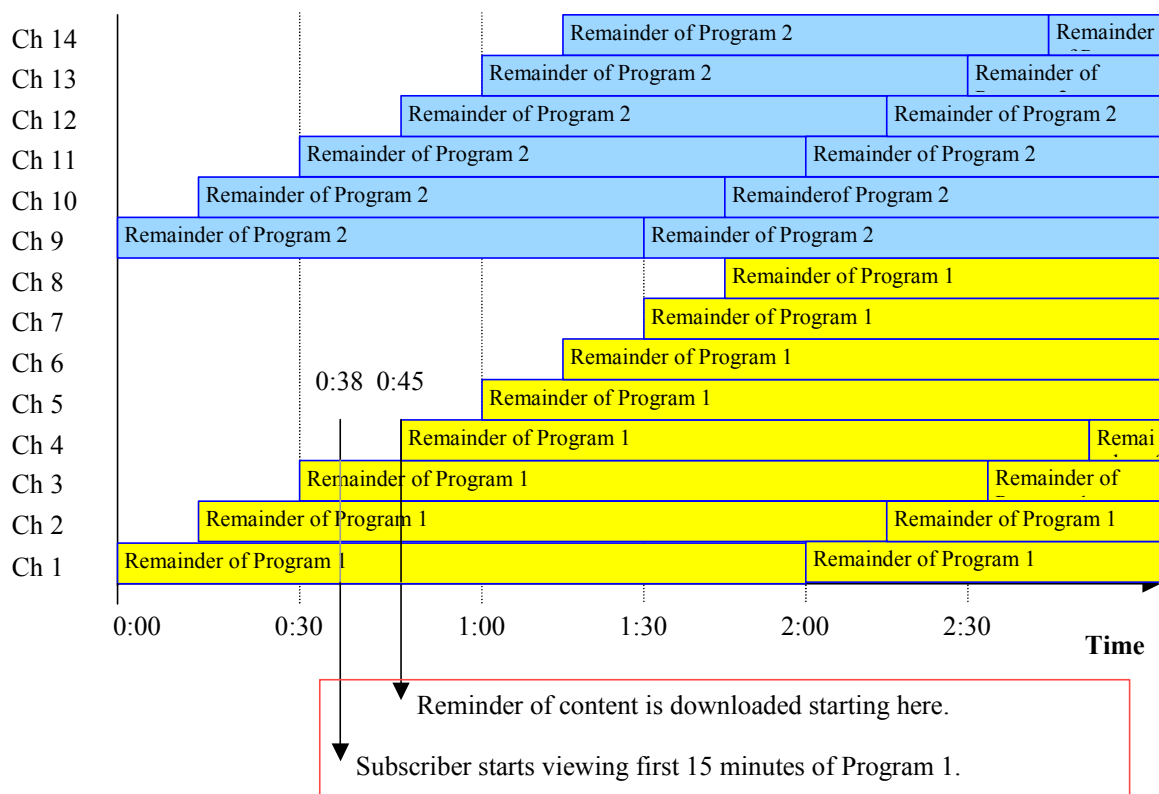
Local Storage of Initial Segments

We present here a more practical method of providing a true VOD service that utilizes local

storage and avoids the problems of the brute-force method while maintaining the advantages of a simpler distribution system. In this method only the first few minutes of each premium program is downloaded to the local set top of subscribing customers. In addition, a carousel of the remainder of the program is continuously broadcast on dedicated channels. In this way the system uses existing broadcast infrastructure with a service similar to existing PPV, but provides true VOD features to the customer.

We will show that this method allows true VOD with all trick-play capabilities without storing every program on the local drive, and without requiring dedicated connections between the media server and each customer. Furthermore, by using technology similar to existing PPV systems, this VOD service can be implemented using existing technology.

Assume that the first 15 minutes of each premium access program is downloaded to the local drive in the set-top. The program snippets could be provided via a single dedicated channel that broadcasts during off-peak times. Software on the



set-top would add previews of new programs as they appear on this dedicated preview channel, and the software would remove content when it becomes outdated.

Because the first 15 minutes of every program is available to the customer, they can receive instant, on-demand access to the premium content and begin viewing at any time, therefore providing true VOD. Customers also have the ability to perform full trick-play of the content since it is stored locally.

To provide the remainder of the premium access program (minus the first 15 minutes), a continuous carousel of the remaining portion of each program is broadcast on digital channels. Multiple versions of each program are broadcast on different channels, with each version delayed from the previous one by 15 minutes. For a two hour and 15 minute program (the example shown for Program 1 in the figure), this would require 8 channels to broadcast each version of the program. Once the customer begins viewing the conditional access content (and has completed the purchase transaction), the set-top begins storing the remainder of the content from the start of the next relevant carousel stream. Because the start time of the remaining content is broadcast every 15 minutes, the customer is guaranteed to start receiving the remainder of the content before they exhaust the 15 minutes already stored on their set-top. To the customer it appears as if the content was immediately available without delay.

An example of the system is shown in the figure. In this example two programs are distributed in carousel fashion across 14 channels. The total length of Program 1 is 2:15 and that of Program 2 is 2:00. In the example, a customer begins viewing Program 1 at 8 minutes after the hour. Since the first 15 minutes are already cached on the local hard drive, the customer is viewing the content directly from the local storage in his set-top box. At 15 minutes after the hour, the set-top starts recording the remainder of Program 1 from channel 4. This is appended to the first 15

minutes already on the hard drive and therefore stores the entire program onto the hard drive.

The system uses a distribution method already used in some deployments for PPV or multi-channel conditional access channels (HBO-1, HBO-2, etc.) where multiple versions of the content are displayed at various time offsets. However, unlike PPV the user does not need to wait until the next start time to begin viewing the content. In addition, the user has full trick play control over his viewing.

The storage requirements for a 15-minute preview are only 0.5 GB per program. The top ten preview programs could therefore be stored using about 10% of the drive capacity (assuming that the hard drive capacity is 40 GB). More importantly the distribution system does not require dedicated direct and persistent connections between the media server and each customer. Instead it uses a carousel system similar to existing PPV systems.

Security issues can be minimized in a number of ways. The content stream is encrypted on the hard drive and requires the resident security system for decoding. Since most security systems include keys with finite lifetimes, the content can only be viewed during a finite time period. Set-top software can also automatically delete the content, including the previews, after a finite time. Since the preview material is more accessible to potential hackers, it could be secured with a separate system such that if the preview programming is compromised it does not allow viewing of the entire program.

Storing previews on the set-top

Advantages

- Simpler network infrastructure requirements
- Uses familiar PPV carousel method.
- Provides PVR services to each customer
- Compatible with existing VOD infrastructure

Disadvantages • Some security issues with locally stored content. • Locks up some percentage of the available local storage
Cost Per Subscriber • \$100 - \$180 + adder for hard drive • Multimillion dollar VOD infrastructure not required

In summary, cable operators are continually searching for products and solutions that can deliver double digit operating cash flow (OCF) growth. Keen Personal Media has developed an elegant solution available through its TV4me PVR technology that leverages the existing cable system's infrastructure, while providing cost effective software and hardware. The scalability and costs associated with each solution are significantly disparate as the proposed methodology offers a 10 to 1 cost advantage over traditional VOD infrastructure costs. .

This solution essentially moves storage to the edge of the network, with the hard drive residing in customers' homes. From a cost perspective, this solution has significant advantages over the required expenses for transport and storage located in the head-end. The primary difference is the need for peak usage allowances, which require an excess of bandwidth and VOD streams. The proposed schema optimizes load balancing of content transmitted within the cable infrastructure providing the most cost-effective solution for cable operators.

This paper is not intended to suggest that local storage is a direct replacement for current VOD systems, however we do suggest that local storage in conjunction with a VOD system could deliver the optimal solution for full VOD services to a wider population while minimizing head-end infrastructure. A hybrid system consisting of deployed client set-tops with integrated PVR storage and some headend VOD capability would be more scalable and offer the operator a higher degree of flexibility than current systems and have a longer lifespan,

ultimately creating a more reasonable solution and potentially greater customer satisfaction with less attrition.

Contact Information:

Keen Personal Media, Inc.

One Morgan

Irvine, CA 92718

Tel: (949) 672-6380

Fax: (949) 672-6211

E-mail: TElliott@KeenPM.com

Bibliography

¹ Adams, Michael. "Open Cable Architecture," Cisco Press, 2000.

INTEGRATING HARD DISK TECHNOLOGY: ENABLING INTERACTIVE TV

Graham Williams
Pace Micro Technology plc

Abstract

The introduction of Hard Drives into the "Home Gateway" space will change the way people watch TV. This paper addresses some of the engineering issues associated with introducing this technology, and also some of the benefits it will bring.

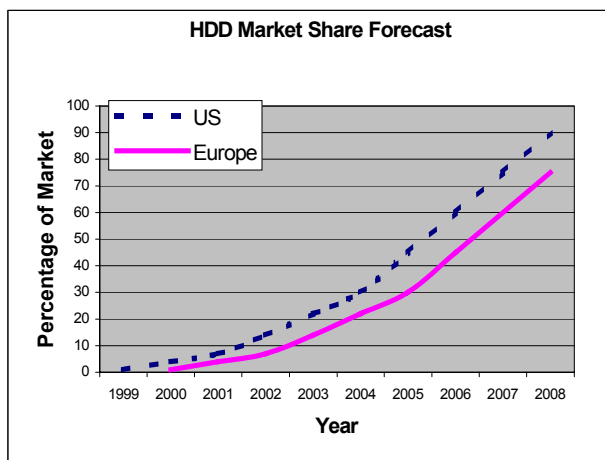
OVERVIEW

Definitions

DVR: This will refer to the basic set of functions that can be associated with the standard VCR today, such as Fast Forward (FF), Rewind (Rew), Pause, Record and Play.

PVR: This will refer to the enhanced personal features such as recording individuals favorite programs, or actors, movies etc.

Pay PVR: This will refer to push services that will earn additional revenue for the MSO, such as downloading the top 5 PPV movies or video games.



Market Forecasts

Market studies predict that by 2005 approximately 40% of the Home Gateway market will be for products with an integrated HDD. For this to happen the hardware and software designers need to work to make the products affordable and to have a perceived market value. There are a number of products that have launched to date that can only be seen as the first in a long range of products and educating the consumer to the benefits of the new technology will be one of the major challenges for the market forecasts to be met.

THE ENVIROMENTAL CHALLENGES

The consumer environment has to be considered as one of the harshest environments for the HDD to live. A consumer will demand high reliability while treating the product in a similar way to any other consumer piece of equipment. The box can regularly be seen in a A/V cabinet with a number of other heat generating pieces of equipment such as VCR's, DVD's and Home Theater Amplifiers, it may also be seen positioned next to or on top of a speaker or even worse a sub woofer and also has to remain totally silent during those quiet points in the latest movie. These issues all need addressing if the product is to be successful.

Heat

Today HDD's are generally designed to operate at a maximum operating temperature of 55 deg C. While this can be seen as acceptable in the PC environment it demands attention in the Home Gateway space. Most consumer products are specified to work in an ambient of up to 45 deg C and are still expected to maintain a mean time to failure of

300,000 to 500,000 hours. On average the ambient within a gateway is between 10 and 15 deg C above the external ambient, hence simple addition shows we instantly hit the maximum operating temperature of the drive. This is being addressed in two ways, firstly as drives become common in the consumer environment drive manufacturers will start to design their drives for this rather than just the PC market. This has started and we will soon see drives with 65deg C max operating temperatures. Secondly today we find that most products require forced cooling of the drive. Fitting a fan, which in turn has its own design problems of noise and reliability, performs this in most cases. Careful design of the airflow in the enclosure to maximize the airflow around the drive gives maximum cooling for minimum fan usage. The fan should also be controlled to be at least software switchable so that its only used as the temperature rises, and preferably speed controlled as this helps to extend the mean time to failure of the fan and also limit the audible noise generated.

Vibration

The design of the mechanical mounting for the drive is one of the most important aspects to consider. It's true to say that a drive that performs well in one product will not necessarily perform well in another. This can be attributed to the mechanical mounting. Vibration due to the disk spinning or the heads seeking, as well as external vibration from speakers or even people walking across the floor can all cause data to be miss read which in turn instigates another read cycle which introduces a time delay. This can be acceptable in the data world but will automatically cause a disruption of the picture, and hence a consumer complaint.

Noise

There are three main sources of noise to consider, rotational noise from the disk spinning, noise generated by the heads seeking, and noise generated by the fan introduced to maintain the operating temperature of the drive. The rotational source can be attributed to the quality of the drive but also the quality of the mounting discussed early. Noise generated by the heads seeking can also be attributed to the drive quality but can be minimized by the software drivers and file management system if designed to minimize the amount of seeks required. The fan noise can be addressed by choosing one with a good quality bearing and by having control over it's speed in software, which keeps the noise generated to the minimum necessary for temperature control.

TECHNICAL CHALLENGES

Analogue vs. Digital

Storing compressed MPEG2 video takes many gigabytes of memory. Recording a single broadcast quality program requires 2Gbytes for every hour (based on an average bit rate of 4.55Mbits/sec for all components of a program). A HDD therefore requires a minimum of 6Gbytes for a 3 hour film and realistically a minimum of 20Gbytes to be useful. Today we see 40Gbytes becoming the standard size this will reach 80Gbytes by the end of 2001 and 150Gbytes by the end of 2002.

Existing stand-alone DVR boxes today take in analogue video and audio and compress the signals to MPEG before they're recorded to the drive. This allows the picture quality or resolution to be controlled by the user in a similar way to VCR's today where the user can select SP (standard play) or LP (long play). This is not possible in the digital world as the video source is already in the MPEG format when transmitted hence the record

resolution is always the same as the broadcast source. This adds complexity to the user interface and the scenario below becomes an issue.

A 2-hour event can take up different amounts of drive space depending on the resolution and hence a 2-hour film on one channel may fit into the remaining drive space whereas a 1 1/2 hour football game may not. Its therefore not possible to directly compare a standalone DVR box today that is specified to record 20 hours of content with an HDD integrated into a product that receives MPEG content directly.

Maximum data rates on HDD's today

Regardless of bus interface, HDD's all fundamentally work in the same way. They're organized as a stack of 1-6 coated metal or glass platters that rotate together. Each platter has 1 or 2 recording surfaces. There is a head assembly with one read /write head for each recording surface and all heads move together. A disk surface is organized as thousands of concentric tracks each divided into a number of individually addressable sectors. Consecutively addressed sectors lie next to each other on the disk, either on the same track on a different surface or on the immediately adjacent track. When a command arrives, the heads seek to the correct track and waits until the first sector rotates under the head (rotational latency). Data transfer then happens to/ from the disk for some consecutive number of sectors. It should be noted that during a data transfer operation the seek rotational latency times are overheads no media data transfer is actually taking place. Once the heads are on the track data transfers happen very quickly with almost no overhead. This media transfer is the ultimate bandwidth limit for HDD's. For small transfer sizes the seek and rotational latency overheads dominate. Therefore to get high sustainable bandwidths you must do large sequential

accesses. This means transferring the data in the maximum block sizes allowed to minimize the overheads.

Today's hard drives have basically two standard interfaces, SCSI and ATA. SCSI is aimed at high-end devices hence costs too much for consumer products, hence ATA is favored. ATA is also commonly known as an IDE interface which is a cheap efficient interface over a 40 pin connector. ATA is both an electrical and command specification. ATA comes in a variety of headline speeds, PIO mode, multiword DMA, UDMA/33 (33MB/s), UDMA/66 (66MB/s) and soon UDMA/100 (ATA100 – 100Mb/s). For streaming content these "headline" rates mean nothing, they're just peak transfer rates to and from the HDD cache across the interface, sustainable bandwidth is limited by the media transfer rate which is today an maximum of approximately 25Mbytes/s or 200Mbit/s. Therefore UDMA/33 interface is sufficient for DVR.

As mentioned the sequential transfer size is very important for streaming video. The current ATA4 specification has a maximum limit of 128Kbytes for a single transfer. ATA5/6 will address this limit and increase it up to 32Mbytes.

As the main market for HDD's today is the PC market they're designed to maintain data integrity. This means that there is a very high upper limit for the time a command will take to complete. Unfortunately this makes it very difficult to schedule real time data transfer for video recording and playback. Currently we must use a sledgehammer approach and switch off almost all error correction. ATA5 and 6 will address this issue and provide methods for to control the maximum time taken for executing a command. This will allow the HDD to dynamically adjust the amount of error correction that it performs

whilst still allowing hard real time scheduling. In either case if a non-correctable error occurs for the file system access (i.e. data not video) the HDD driver software must take actions to retry that block again. This is an example of how the consumer DVR market is beginning to steer the HDD manufacturers.

The cost and complexity of legacy analogue support

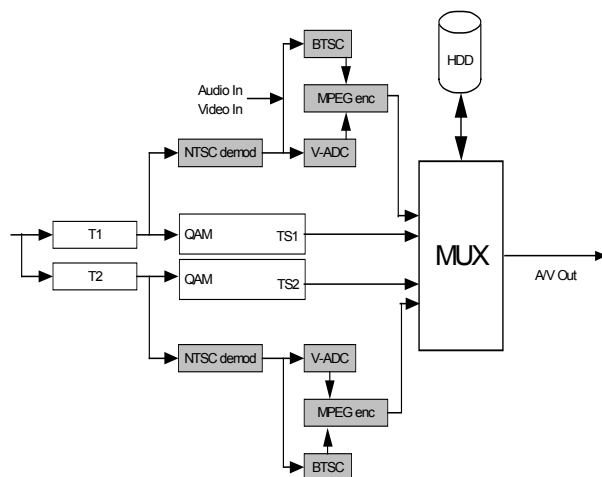


Figure 1

If we were to list the most wanted features for a DVR product the so-called *trick modes* would be near the top as would to watch one channel while recording from another. This should be seamless to the consumer therefore there needs to be two full paths for the video delivery, as shown in fig 1. The analogue channels need to be digitized and MPEG encoded before being stored to the drive, which means supporting two real time MPEG encoders, 2 BTSC decoders, 2 video digitizers and 2 full band tuners capable of receiving both analogue and digital channels. The shaded sections in fig 1 represent this. If there were any one thing that would simplify the hardware architecture and improve the overall experience of an integrated digital cable DVR product it would have to be a digital only solution. This may mean carrying all analogue

channels in the digital domain also which brings additional head end cost, complexity and inefficient use of cable bandwidth but the overall improvement may be worth the investment.

Network vs. Local storage

VOD could be regarded as a form of DVR and in many ways would appear transparent to the consumer. It can support FF, Rew, Pause and Play, however some of these have more bandwidth implications to the system. For instance assume that you had paid to watch a film and halfway through you needed to pause for an hour for some reason. If the storage were local in the box then no problem, assuming that there's space on the drive the film would be recorded and would start to play from the pause point after the hour.

In a VOD (networked storage) system this would mean that the session would need to be extended by an hour which in turn would impact all the assumptions made for the VOD system bandwidth calculations, plus as DVR becomes more and more heavily deployed this loading would probably increase at a greater rate as the consumer became more familiar with the DVR features. This session could be further extended if the user decided to rewind the film halfway through so that a second person could watch it. There are some policy decisions that need to be taken by the MSO to control this or the sessions can become longer and longer causing serious bandwidth issues.

The answer may lie in a hybrid model where the local storage could buffer the VOD content to limit the session to the length of the film, and as a next step it could be used to reduce the VOD session further by allowing the content to be downloaded at a much faster rate, up to 38Mbits/s which could mean a 2 hour film could be downloaded in approximately 15 mins. The hybrid model

would also allow for the amount of storage to be more readily upgraded on the network side and hence would extend the life of the local storage product. There will no doubt be a sweet spot that justifies the business model that will allow mass deployment of HDD enabled boxes but that can be future proofed by network upgrades.

FUNCTIONAL CHALLENGES

Moving from DVR to PVR

To gain the advantages of PVR there's a requirement for enhanced or extended guide data (metadata) and a software module sometimes referred to as a preference engine that filters the guide data to find the user or multiple users preferences such as, all John Wayne movies, or all new episodes of Friends. This then extends down to the file management system to manage the drive. The user interface then has to deal with the remaining drive space and how much drive space each user may have access to for instance do I record football for dad or Friends for mum? The quality of the PVR experience will heavily depend on the quality of the additional guide information and the preference engine, which in turn will drive the data providers to create this enhanced data.

Moving from PVR to PayPVR

PayPVR brings the additional challenge of how to pay for the drive in the box. While DVR and PVR both give the user an enhanced experience neither justify significant additional revenue for the MSO. Only when services can be downloaded or "pushed" to the drive can additional revenues be generated such as downloading the top 5 PPV movies to the drive to help smooth out the peak demands on a VOD, home shopping catalogues or video games. The addition of PayPVR adds conditional access issues to resolve such as

does the user get billed when the movie is downloaded or when its watched. This is an important issue as today most conditional access systems were designed before this became a reality.

Archiving

There will always be the desire by the consumer to archive the content on the drive to some other form of storage. This today would most likely be VHS Tape but in the future may be Digital-Tape, DVD or even a second HDD. While the content is stored in its analogue form this is deemed acceptable to the MPAA but once the content is stored in its original digital format then the MPAA will require a form of copy protection, this is commonly referred to as Digital Rights Management or DRM. Most silicon providers are now incorporating a mechanism to allow for the content to be encrypted before being stored on the HDD. While there's no official standard today this must be addressed or the advantages of HDD will not reach their full potential.

IN SUMMARY

The integration of HDD's into consumer products is inevitable and the technical issues discussed above have or are being addressed and in the future we will see a wider variety of consumer products with HDD's becoming available. This will drive down the cost of the products, which in turn will drive the deployment and encourage new interactive applications that can take advantage of the technology.

IP NETWORKING IN HFC BROADBAND NETWORK

Oleh J. Sniezko, Vice President of Broadband Engineering
Xiaolin Lu, Vice President of Strategic Engineering and IP Networking
Luisa Murcia, Vice President, Data Services Technical Operations

AT&T Broadband, 188 Inverness Drive W.
Denver, CO 80112

Abstract

This paper describes the IP networking trends, the existing solutions, and possible evolutionary implementation of an integrated IP platform over hybrid fiber-coax (HFC) broadband telecommunication networks. Some of these implementations are analyzed based on the trials conducted and planned by AT&T Broadband.

In the first phase of these trials, the IP network is designed to enable AT&T Broadband ChoiceSM (ABC) for our customers. It provides them with a choice of multiple ISPs while at the same time enabling IP telephony supported by a broadband telephony interface (BTI).

The next phase of the planned trials will implement an end-to-end voice over IP (VoIP) solution. The ABC networking will be continuously optimized as more advanced routing equipment becomes available.

The bulk of this paper concentrates on those technical aspects of the field test that relate to networking, design and performance. The authors list and describe the existing equipment used in the trial and analyze the challenges to bringing the service reliability and quality to a level equal to or better than is being provided by traditional telecommunications networks.

INTRODUCTION

Today, HFC access networks deliver a multitude of services using several, very different multiplexing techniques for downstream signals and several different media access protocols in the upstream path. Similarly, on the network side, signals for different services are distributed over disparate transport systems, and interface with several different backbone networks. This situation exists mostly due to the historical evolution of telecommunications, data and entertainment service delivery systems.

The technological progress in the recent years allows for integration of most of these systems. Operators can now begin the transformation of their metropolitan and HFC networks towards packet-based service delivery networks for most of the interactive services. The metro network can be fully integrated and the multiple service platforms (MSPs) are becoming readily available from traditional transport vendors and from new vendors. These platforms are suitable for distribution of video, data and telephony signals. They often integrate layer 2 and layer 3 functionality. Together with the positive cost trends in video encoding technologies and in MPEG flow processing equipment, signals for all services, including broadcast analog video, can be distributed over these integrated platforms.

Similar integration on the HFC access side is technologically possible but

not currently practical. The reasons for this are: the traditional consumer electronic equipment characteristics, some regulatory requirements and cost implications of the integrated approach. Most of the video services, including analog and digital broadcast video and VOD as well as analog and digital ad inserts are best distributed over a SCM network in their analog form or as MPEG signals. Most of the other services including telephony, data, iTV, signaling information, and the like can be cost-effectively transported over an integrated IP DOCSIS platform.

The platform integration can be accomplished in a one-time transformation or in a staged manner. In a staged transformation approach, the IP networking can be initially implemented in primary headends where different services, with the exception of broadcast video and MPEG VOD signals with local ads, could be converted for distribution in a packetized form over a DOCSIS-based platform to our customers. On the network side, the services could be initially separated to interface with different legacy transport platforms. From this initial point, the evolution could progress in two different directions. Pure packet-based systems could be driven deeper into the HFC access network with the concept of distributed CMTS. This approach would allow for simplification of networking in the access section of our systems. Only the last mile would remain optimized for RF SCM coaxial delivery. There, broadcast analog signals would be carried in their original format, and digital broadcast video and VOD signals would be carried in MPEG format, QAM modulated onto RF carriers (with digital ads spliced in). Most of the other signals would be carried in their packetized format over DOCSIS to the customer terminal devices.

In parallel, the packet-based networking would expand in metro networks

to evolve into an end-to-end IP solution for all services, including traditional video services and telecommunications services. At the same time, some interfaces would be required to account for the fact that many of these services are carried on legacy systems such as PSTNs.

OBJECTIVES

Future IP networking focuses on the design and development of an AT&T owned and operated IP network, and the design and development of the client products that create value for our customers. The design provides our high-speed Internet access customers with a choice of multiple ISPs. The trial of this new infrastructure allows for testing how this network plugs into AT&T's current high-speed data backbone. It also supports IP telephony (IPT) on the access side. Moving forward, this IP network will support not only high-speed data (residential and commercial), but also iTV and other interactive services.

There are two primary differences between this new infrastructure and the conventional IP infrastructure:

1. Ability to efficiently provide AT&T Broadband's customers with convenient access to multiple ISPs and their services;
2. Ability to support multiple applications.

To test these features, AT&T Broadband asked several ISPs to participate in the first phase of the trials. The trial is taking place in Boulder, Colorado in the area with 25,000 homes passed. The following ISPs are currently participating:

- EarthLink,
- Juno,
- Worldnet, and
- Excite@Home.

This group may be joined by RMI.

AT&T Broadband has secured 400 volunteer customers of which 333 are

already participating. Each of the multiple ISPs introduced tiered services. The table below lists the tiers of service for each ISP currently participating in the trial.

Table 1: HSD Service Tiering

ISP	Tier 1 # of Users @ 1.5 Mbps	Tier 2 # of Users @ 300 kbps	Tier 3 # of Users @ 128 kbps
EarthLink	21	1	3
Juno	47		
Worldnet	132	51	13
E@H	65		

The tiers presented in Table 1 reflect downstream capacity. The upstream capacity is held constant @ 128 kbps for all tiers.

PLATFORM DESCRIPTION

Generic Requirements

Supporting the infrastructure of the AT&T Broadband ChoiceSM (ABC) involves capabilities in three key areas of the infrastructure:

- a user-friendly interface that can run on the end user's PC to provide a convenient method of selecting services and ISPs;
- policy routing capabilities in the underlying routing network so that traffic destined for different ISPs can be sorted and routed accordingly; and
- a Service Activation System (SAS) for service provisioning and activation in the AT&T Broadband ChoiceSM (ABC) environment, capable of migration to an integrated provisioning platform for multiple-service offering.

GENERIC DESCRIPTION OF INFRASTRUCTURE

Integrated Platform in Access Network

The HFC access infrastructure was designed to meet DOCSIS 1.1 requirements, but the access platform was based on equipment compliant with DOCSIS 1.0. To support high-speed data delivery and VoIP over DOCSIS, the equipment was modified to meet the objectives of the DOCSIS 1.1 standard before the CableLabs certification commenced or to be capable of migrating to DOCSIS 1.1 standard. Both elements: CMTSs and BTIs, were deployed. Many other parameters related to VoIP in the HFC environment and covered by the evolving PacketCable standard had not been certified.

High-Speed Internet Access

On the network side of the CMTSs in primary hubs, the traffic for high-speed Internet access is separated from the telephony traffic. The HSD traffic is routed to different ISPs based on customer choice.

Voice over IP

After being separated from the IP data traffic, the IP voice traffic is routed via traditional gateway (IPDT for IP digital terminal) to Class 5 switches. This configuration on the network side does not differ from the existing approach; the telephony traffic is treated separately. A generic diagram of this network is presented in Figure 1. Additional expansion will be implemented to improve network reliability by creating a redundant point of connection (PoC) as presented in Figure 2.

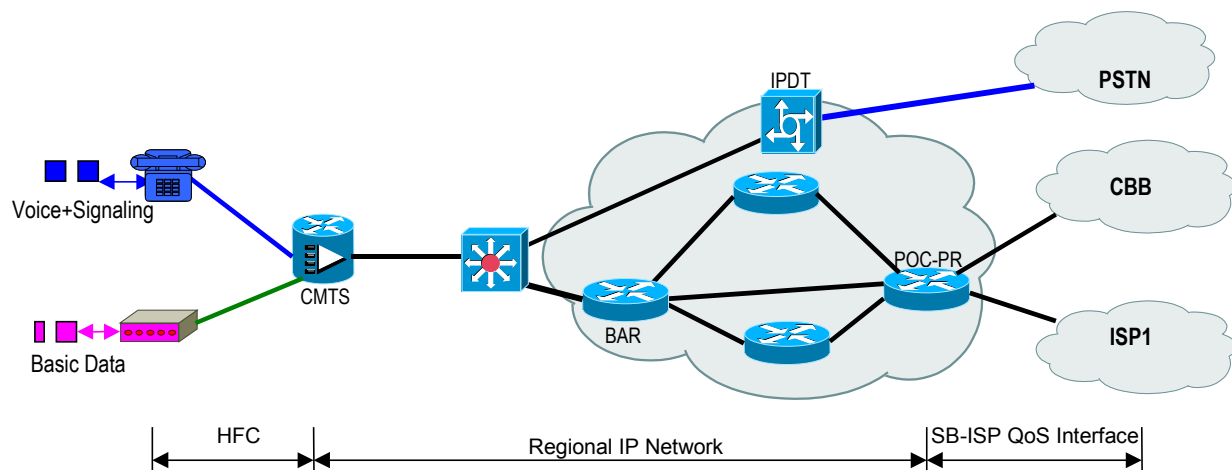


Figure 1: Generic Architecture for Integrated, Multi-ISP and Multi-Application, Broadband HFC, DOCSIS-Based Platform

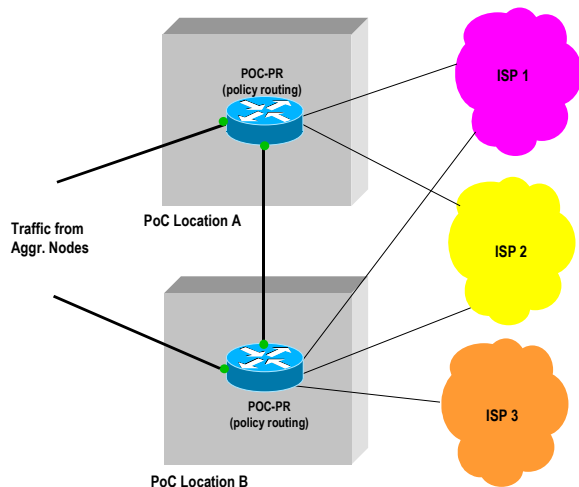


Figure 2: Redundant Configuration of PoCs for Higher Reliability

MULTIPLE ISP NETWORKING

Layered Broadband IP Infrastructure

The multiple-ISP network comprises two main sub-networks: regional broadband access infrastructure and intra-regional backbone. ISPs can interface directly with the regional network (local ISPs) or via inter-regional backbone. A simplified configuration of this layered approach is presented in Figure 3.

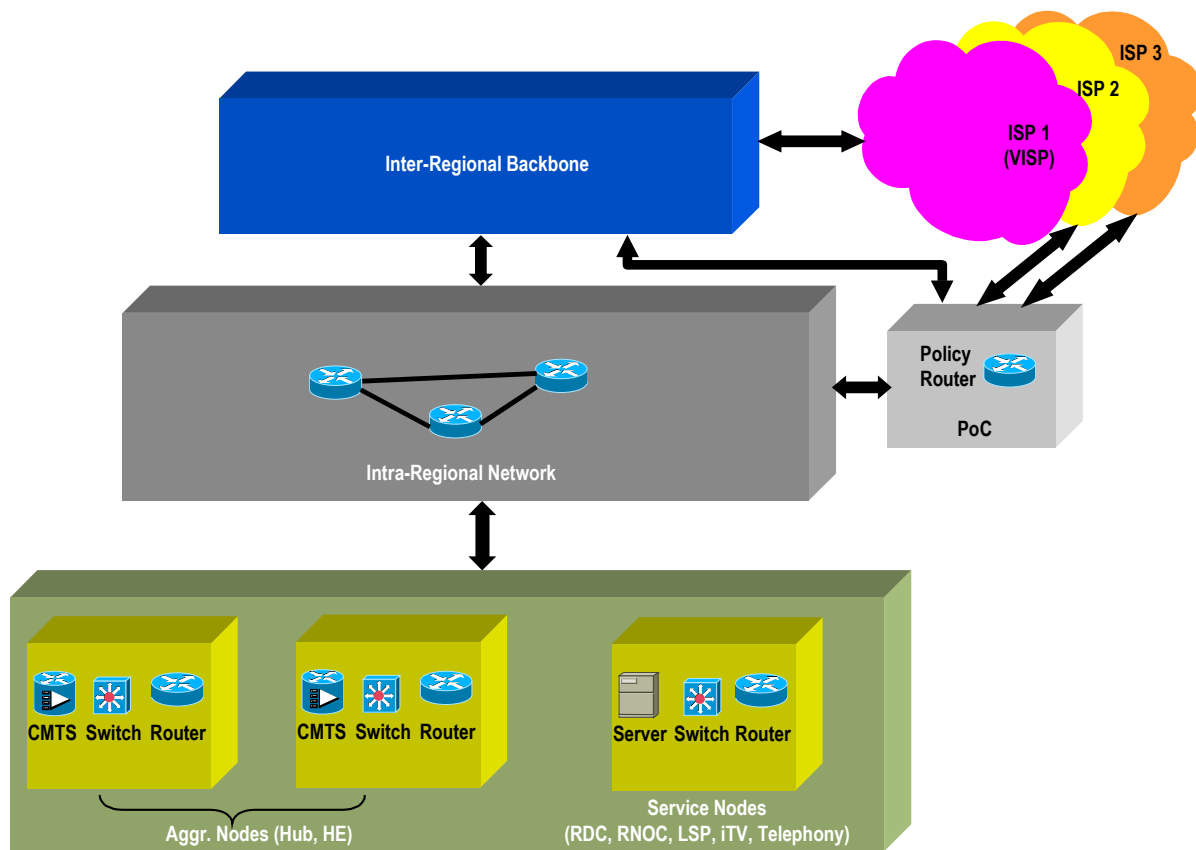


Figure 3: Regional and Inter-Regional IP Infrastructure

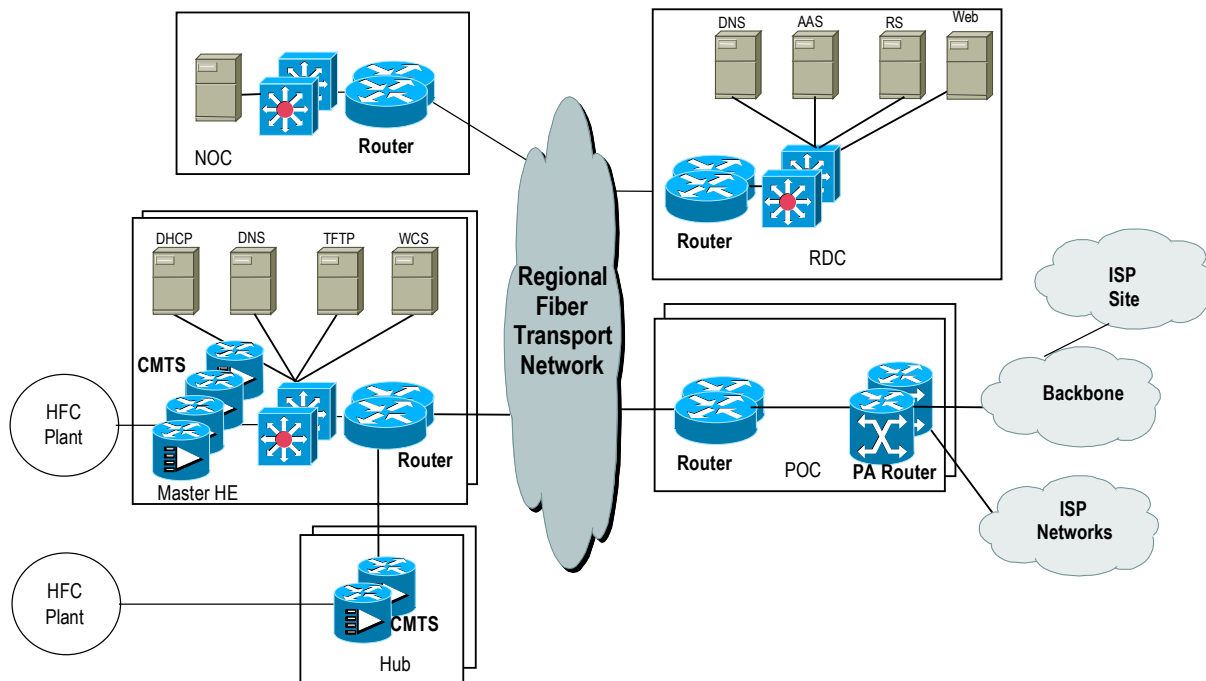


Figure 4: Regional Broadband IP Infrastructure

Regional Broadband Access

Figure 4 depicts the components of the regional infrastructure. The regional network:

- aggregates end-user traffic in aggregation nodes located in hubs and headend; and
- provides local services such as local ISP connections, regional registration, and regional NOC through service nodes.

Customer base projection and network/traffic engineering determine geographical size and boundary of a region.

Inter-Regional Backbone

The inter-regional backbone allows for connecting remote ISPs (virtual ISPs) to the regional network. It also provides a means of distributing content from a central location to the regions. It is used by applications/services requiring inter-regional connectivity. It also provides connectivity to a centralized NOC

Policy Routing

In the ABC environment, one of the major tasks is to route traffic to the ISP of choice in an optimal way to:

- accurately determine the allocation of resources and billing per an ISP; and
- to avoid congestion in the transport systems (including inter-regional backbone).

At the time when the first phase of the trial commenced, source address-based policy routing was developed to meet the requirements of the ABC environment. In this approach, a PC is assigned an address from the address block of the ISP of choice. The packets going out of the PC are forwarded to the ISP based on the source address (policy routing). The incoming packets are routed in a conventional way to the client PC based on its address. This routing method is reflected in Figure 5.

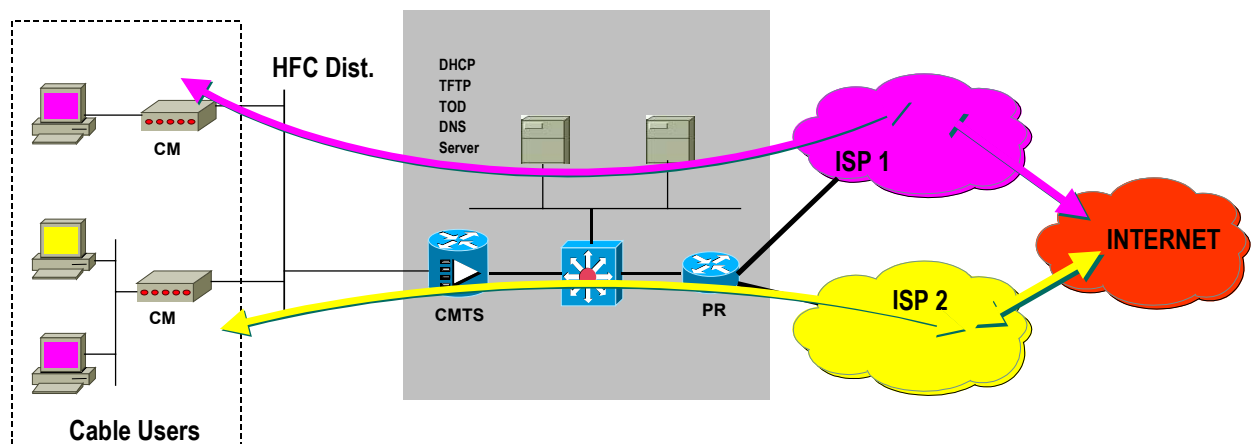


Figure 5: Operation of Source Address-Based Policy Routing

This method of policy routing was supported during the trial by a single policy router (Cisco 7513). The choice of this

policy routing was based on technology availability and test results. Some limitations existing at the beginning of the trial were

later eliminated by improvements in the policy router. The trial results confirmed that this method of policy routing is effective. However, it revealed some shortcomings such as:

- limited number of ports available on a single router;
- no redundancy provided for ISP connections into the region;
- no load balancing provided for traffic to an ISP;
- policy decisions need to be provisioned (and re-configured) centrally on the single policy router; and
- requires a very high performance policy router to support the large number of policy filtering decisions that may be required in a region.

Many of these shortcomings can be addressed in a redundant architecture of PoCs (refer to Figure 2). In this configuration, policy routing is applied to incoming packets

on all the marked interfaces of PoC-PRs. These PoC-PRs may be collocated. The incoming traffic from the region is load balanced into the PoC-PRs. ISPs with connectivity to both PoC-PRs benefit from load balancing automatically.

With the progress made in routing protocols, a different method of policy routing can be adopted. One method that is being considered for the next phase of the trials is MPLS-based policy routing. The application of MPLS (multiprotocol label switching) and LDP (label distribution protocol) will allow for many other features and will provide sufficient improvement in provision of QoS and CoS. An MPLS-based routing table can be constructed for each ISP. This table would span all PoC locations to which the ISP is connected. Under this scenario, ISPs can connect to as many of the PoC locations as they choose.

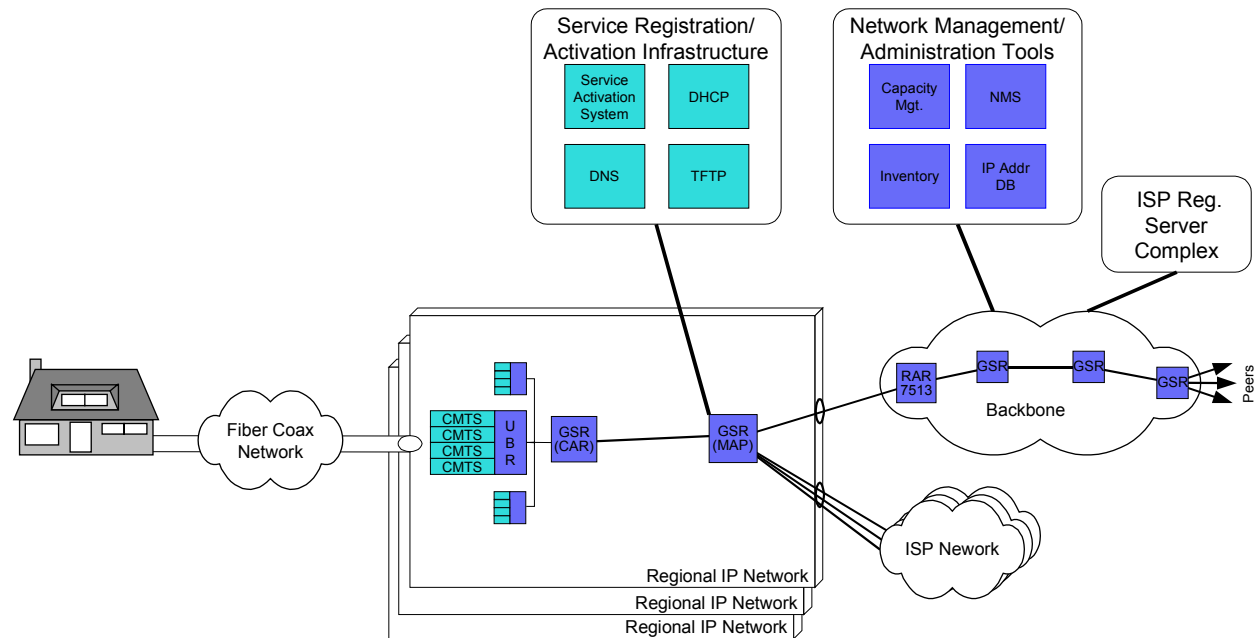


Figure 6: Network Management and SAS Infrastructure

Network Management

The service provisioning and activation as well as network management in

the multiple ISP environment becomes critical. The objective of system scalability and adaptability to the new services required that a significant effort be dedicated to the

development of a suitable platform. The configuration of the network management system is presented in Figure 6

In addition to these requirements, the issue of management of the IP subscriber addresses allocated to ISPs had to be resolved. This issue was especially critical for source address-based policy routing. A detailed address allocation plan was developed. It allows for re-use of Net 10-address space between regions but this feature is not necessary. In the newly built networks, the address allocation does not present a major challenge. However, during migration of the existing networks to the ABC service, a careful evaluation of Net 10 usage is essential. Some address re-assignment may be necessary but can be avoided by careful planning. Address management tools, SAS (DHCP) address usage reporting and automation of address management are essential components of successful address management in the ABC environment.

VOICE OVER IP OVER DOCSIS (IP TELEPHONY)

Generic

The same integrated broadband IP platform was used to provide telephony services. On the access side, BTIs (broadband telephony terminals) were used at customer premises. The calls were terminated into CMTSs located at the aggregation nodes. At the time the trial commenced, DOCSIS 1.1 and PacketCable compliant equipment was not available. Moreover, call management servers (CMSs) and soft switches were either not available or could not provide calling features commonly considered as basic. To provide these features, the calls originating in IP DOCSIS network were routed to Class 5 switch via IP digital terminals (IPDTs), known also as

smart gateways, developed for this application to provide fully compliant GR-303 interface. This configuration is presented in Figure 1.

Bandwidth Efficiency and Traffic Models

AT&T Broadband elected to mix data and telephony traffic on shared downstream and upstream channels for two reasons:

1. more efficient bandwidth use, and
2. possibility of simultaneous use of BTIs and CMTSs for digital telephony and data communication (efficient use of equipment).

Statistical data about today's traffic parameters and user behavior were collected to determine the capacity (i.e., a number of users) supported by a single downstream and six upstream (3.2 MHz QPSK) channels (a typical configuration of a CMTS). The methodology used converted a telephony user into an equivalent number of data users and calculated a cumulative number of composite users based on an assumed total number of telephony and data users. With an assumption of low initial telephony penetration and heavy usage of data in the downstream path, it was determined that a single CMTS can support approximately 2,500 HHP and the limiting capacity is the capacity of the downstream channel (highly asymmetrical data traffic at low telephony penetration).

The calculation was performed for the particular trial application. The model, however, is valid for generic applications with adjusted traffic statistics dependent on:

- user behavior for different service types,
- service type,
- compression ratios for IP telephony,
- penetration levels for different services,
- downstream and upstream channel capacity (modulation level on

downstream channel, symbol rate and modulation level on upstream channel),

- CMTS configuration,
- rate limiting policy, and others.

Performance

Nine of the IP telephony loops were tested for 10 loop and data throughput (dial up modems) parameters. The test results were compared against analog loop performance and DT HFC loop performance. The results confirmed that the CMTS/BTI based IP telephony quality below the quality provided by HFC digital telephony. This is understandable in light of the fact that neither BTIs nor CMTSs were tested for compliance with DOCSIS 1.1 and PacketCable standards and BTIs were among the first of this type of terminal.

The major deficiencies were experienced in round trip delays, impulse hits and phase hits. These deficiencies affected a cumulative measure of GoS (grade of service). Moreover, current BTIs were not compatible with most of the dial-up modems. The effective data throughput for the dial-up modem compatible with BTIs was much lower than that for HFC digital telephony NIUs. The test also showed lower than required ringing voltage.

The test results were communicated to the equipment vendors. They were also used in preparation of an RFI for BTIs and IP telephony CMTSs.

QUALITY OF SERVICE (QOS)

The trial IP broadband network was used to provide two different services over shared channels. Moreover, ISP services were offered at three different levels of guaranteed downstream throughput.

Therefore, at the minimum, QoS functions had to enable:

1. Rate limiting, and
2. Priority handling.

Implementation of these two functions involves control and configuration of HFC and IP parameters. These configurations are presented in Figure 7. DOCSIS 1.0 allows for configuration of class of service (CoS) parameters:

- Max downstream rate: enforced per CM by CMTS on a packet basis,
- Max upstream rate: enforced per SID by CMTS via time-slot scheduling,
- Min guaranteed upstream rate: enforced per SID by CMTS via CIR type mechanism,
- Upstream channel priority: used by CMTS when scheduling time-slot grants per SID.

An enforcement of the control mechanisms for these parameters was assigned to the vendors participating in the trials. To allow prioritization of telephony traffic, DOCSIS 1.0+ with some additional expansions was used. It provided:

- unsolicited grant service (UGS), and
- ToS overwrite

In the IP part of the network, a DiffServ model, based on ToS (type of service bits) was applied to enable different priority queues.

Three different tiers of high-speed Internet access service were provided, all of them under Best Effort. The CoS was defined in the CM configuration file. Voice service was provisioned under UGS with a minimum upstream for voice signaling traffic. The ToS bits were reset by CMTS to DiffServ Code Point EF (Express Forwarding) signaling.

This implementation allowed to provide all QoS functionality required for the trial. The task was simplified by the fact that

the IP voice traffic was terminated at the IPDTs.

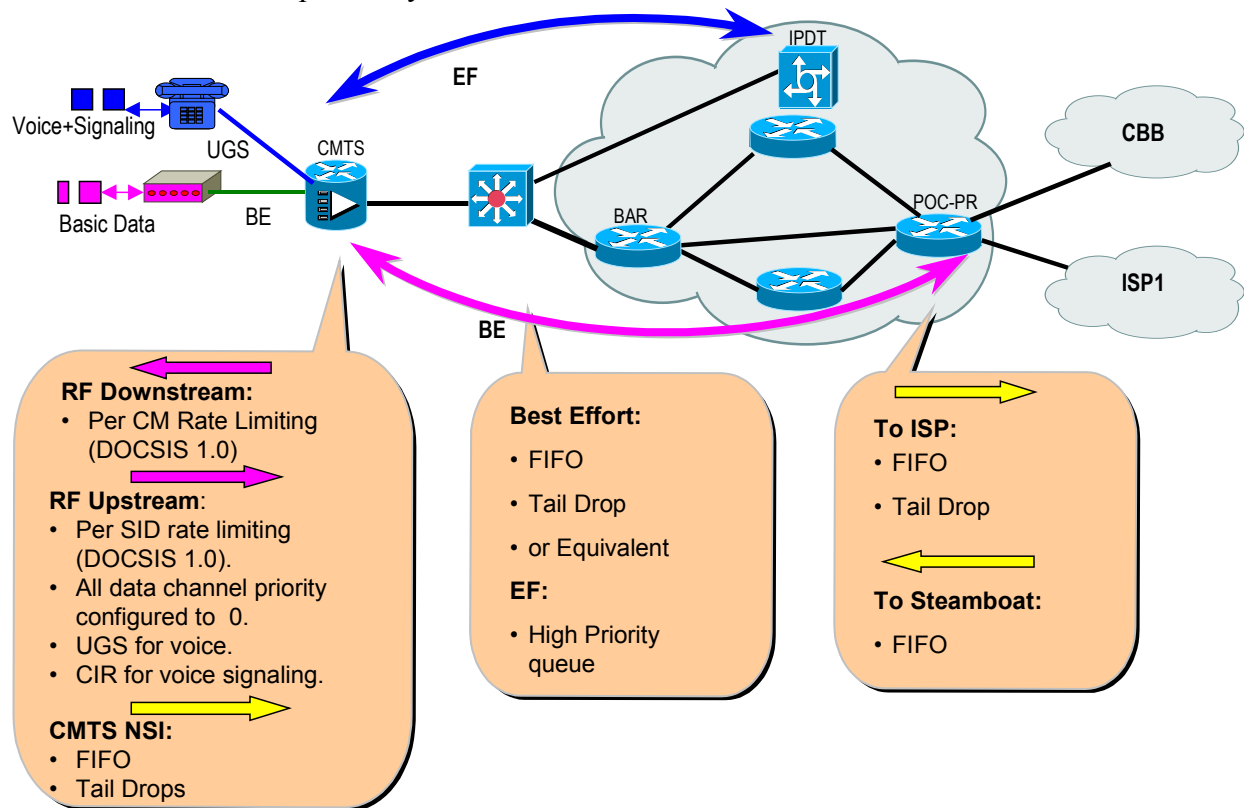


Figure 7: Quality of Service Functionality in IP/DOCSIS HFC and IP Networks

CONCLUSIONS

The first phase of the ABC trial has met its objectives. It provided a unique experience to our high-speed Internet access customers. The new network allowed for implementation of a service model where:

1. Customers can:
 - have a relationship with a selected ISP or ISPs;
 - upgrade the service, change an ISP or any other feature;
2. ISPs can:
 - offer significantly improved service;
 - provide regional or national service;
 - develop applications based on the demand;
3. MSOs can:
 - respond to the marketplace;

- create a specific value proposition to a particular customer;
- maximize value delivered to their customers.

The trial also showed that an additional effort has to be expended to match requirements for IP digital loop performance and make them comparable to those achieved in HFC CBR digital telephone loops. Standard compliant equipment should provide significant improvement in these parameters.

The IP networking required more engineering effort than the setting up of IP telephony. This is mostly due to the fact that IP telephony was contained to the HFC network and the section of the IP network between the

CMTSs and IPDTs. An end-to-end VoIP solution would require a significant amount of effort to implement, even if it complies with the industry standards.

Next phases of the IP network trials will most likely be conducted with end-to-end IP telephony systems and possibly include other IP-centric services. The trial area will pass approximately 100,000 homes. This phase is scheduled in late 2002 but it will depend on the equipment availability.

ACKNOWLEDGMENT

The authors want to acknowledge many of their colleagues from AT&T Broadband and other AT&T organizations, who contributed significantly to the entire project. Specifically, we want to thank Samir Saad, Han Nguyen, Tayfun Cataltepe, Steve Bulka, Pete Zarrelli, Brian Field, Julie Regalado, Cameron Gough, and Allen Stone, but we are certain that many more deserve the acknowledgment.

AUTHORS' CONTACT INFORMATION

Xiaolin Lu

**Vice President of Strategic Engineering
and IP Networking**

Tel: (303) 858-5575

Fax: (303) 858-3181

lu.xiaolin@broadband.att.com

Oleh J. Sniezko

Vice President of Broadband Engineering

Tel: (303) 858-5027

Fax: (303) 858-3181

sniezko.oleh@broadband.att.com

Luisa Murcia

**Vice President, Data Services Technical
Operations**

Tel: (303) 858-3553

Fax: (303) 858-3181

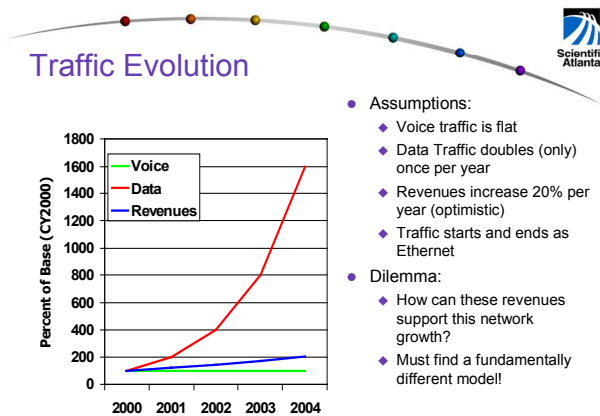
murcia.luisa@broadband.att.com

All authors are located at AT&T Broadband,
188 Inverness
Englewood, CO 80112

IP TRANSPORT TECHNOLOGY IN THE CABLE BACKBONE

Gregory K. Hardy
Scientific-Atlanta, Inc.

Cable operators who have successfully introduced cable modem, voice, and other interactive services to residential subscribers, as well as those penetrating the lucrative business-to-business [B2B] market are beginning to drive a very different backbone traffic pattern, both in volume and composition. We are forecasting a continued growth in IP data as a percentage of the total primary ring traffic. In order to prepare, the cable operator must shift his attention toward the upgrade of the transport network.



Transport Backbone Requirements

For a cable operator and their CLEC affiliate to be competitively positioned, the transport network must meet certain requirements.

Scalability - Ability to cost effectively scale to multiples of tens of Gigabits [Gbps] of transport capacity. As voice, video and data services grow the network should expand on a pay as you go basis.

Low Cost per Gbps – includes initial costs and operational costs. Remember,

bandwidth growth outpaces service revenue growth so we need a new model.

Support of Voice, Video, Data Services – The ability to transport analog/digital video, circuit switched voice, and IP data using a wide variety of interfaces over a single integrated network. It is our belief that IP based traffic will soon comprise the majority, and must be specifically planned for.

Delivery of Next Generation Differentiated Services – The network must support service creation through extensive software Quality of Service [QoS] control, and subsequent monitoring and billing of these services through Service Level Agreements (SLAs). This is particularly important in the B2B market place.

Network Resiliency – 99.999% network uptime, redundant hardware and switched ring or mesh topologies.

Simple but Powerful Network Management - In addition to the critical monitoring and fault isolation capabilities, the ability to simply “point and click” provision the network is very important.

Today's Network Architectures

In touring any major cable headend facility today, you are likely to see one of the following architectures:

1. Multiple networks
2. Metro DWDM [Dense Wave Division Multiplexing]
3. SONET multi-service provisioning platforms [MSPPs]

Most often found is the system with multiple networks simultaneously operating in the cable backbone, each over a dedicated fiber. Distances allowing, you may find analog/digital video on a SONET or proprietary digital network, cable modem likely over a router network, voice and other data transported on ATM. As service traffic grows, a second or third network is often replicated; expect this first in cable modem services as additional CMTS are required. This multiple network approach quickly consumes a majority of available dark fibers and will proliferate facilities with multiple transport platforms.

The introduction and deployment of metro DWDM technology allows the cable operator to reclaim many fibers for future use. Though, two problems still remain. Because metro DWDM platforms typically map a given stream onto its own wavelength, they do not provide a means for efficient flexible aggregation of lower rate [usually electrical] streams in the DWDM platform. Secondly, a transport box per service still remains; each device with its own set of protocols, element management and staffing, and support requirements for the network. The DWDM solution is still complex and expensive.

SONET multiplexers are often utilized to electrically aggregate low data rate services like 100 BT or OC-3 ATM into higher bandwidths, typically OC-48, for transport. Current technologies recently released are based on new mappings of circuit-switched payloads into SONET frames. These products, often referred to as SONET MSPP's [multi-service provisioning platforms], map traffic into different size circuit pipes based on multiples of SONET STS-1. Also available are numerous ATM and VP networks based on ATM over SONET technology. Good news is, this aggregation approach to transport simplifies the cable backbone, particularly in constant

bit [CBR] services that fit neatly into standard SONET pipes. The bad news is, SONET MSPP's are less bandwidth efficient when transporting variable bit [VBR] IP based traffic and that's where all of your growth is projected to be.

A New Technology

There is a new aggregation technology that will increase transport bandwidth efficiency and utilization while still offering the desirable benefits of SONET's protection switching and ATM's QoS. This new technology Resilient Packet Ring [RPR] is currently an IEEE 802.17 working group with expectations to emerge as a worldwide standard in early 2002. RPR combines the best of SONET, Gigabit Ethernet, and DWDM technologies into a box that optimizes IP transport in the cable backbone.

Scientific-Atlanta and our technology partner Luminous Networks are active members of IEEE 802.17 and have submitted our version of RPR to the group called Resilient Packet Transport [RPT], for inclusion in the final standard.

RPT Technology Blocks

RPT technology is centered around four major innovations.

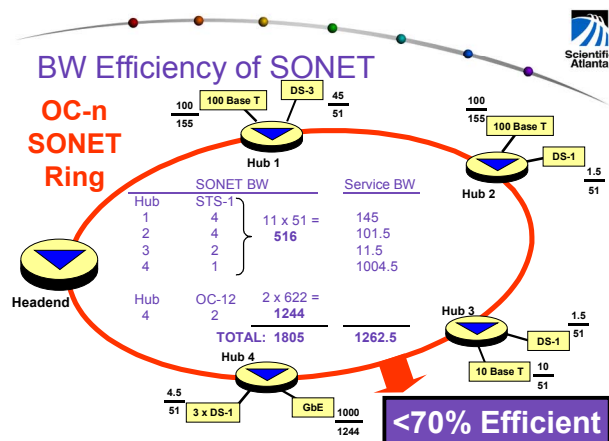
1. RPT utilizes Gigabit Ethernet over Fiber, a packet-switched optical transport layer, optimized for efficient IP transport. This transport layer is similar to an Nx Gigabit Ethernet physical layer, with the addition of robust service mechanisms like performance monitoring, network synchronization, and control packets. RPT achieves SONET-level robustness at Gigabit Ethernet costs.

RPT's physical layer can be single wavelength or DWDM and is decoupled from the logical topology of the network.

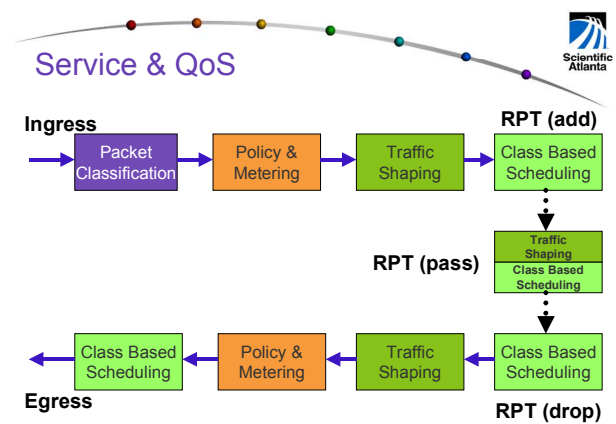
This enables the cable operator to configure point-to-point, rings, and mesh networks.

2. RPT enables full system redundancy and 50 ms span restoration switching. RPT also allows up to 200% of the ring bandwidth [work and protection] to be over provisioned with traffic when there is no span break. Typically, SONET requires the protection ring to be in reserve for protection switching events.

3. All traffic is statistically multiplexed in Layer 2 packets at point of ingress. This provides optimum efficiency in the transport of bursty IP traffic that otherwise would be limited by preconfigured SONET pipes. RPT yields a 30 to 50 percent increase in effective transport capacity over comparable circuit switched networks operating with similar nominal bandwidth links.



4. RPT provides extensive QoS support, including policing, shaping, class-based queuing, and flexible scheduling employing strict priorities and algorithms. RPT also supports minimum latency express paths for packets passing thru [but not dropping] at a hub. RPT overcomes the best effort IP QoS limitation found in Gigabit Ethernet to offer a strong alternative to ATM.



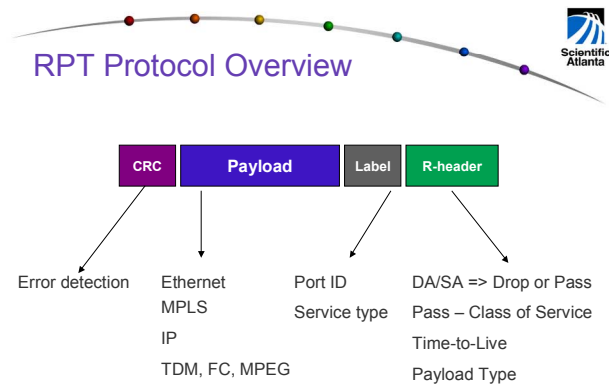
Implementation Benefits of RPT

In implementation, RPT will allow the cable operator to effectively deploy new services, simplify network design, and reduce operation expenditures.

RPT eliminates adaptation to connection-oriented protocols [like SONET and ATM] and recovers some lost overhead. These layers are replaced by a thinner and more flexible RPT and MPLS stack capable of transporting IP services, TDM voice circuits, and video.



The RPT packet includes a 6 bytes RPT header and 4 bytes MPLS and contains all the necessary information to move the packet through the network. This is a very small tax relative to SONET and ATM for similar functionality.

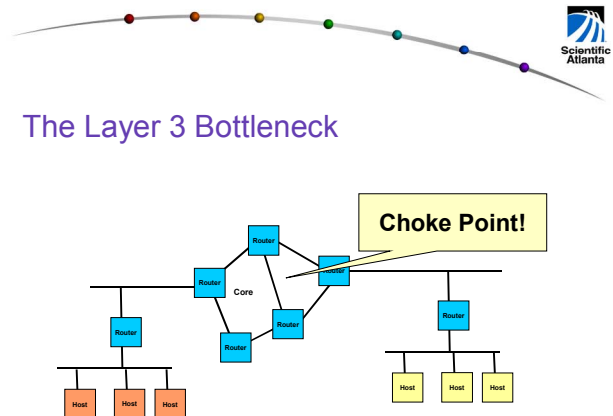


Central to RPT is the ability to enable switching and transport at the packet level. Packet level processing yields infinite granularity and maximizes transport bandwidth efficiency. Eliminating the need to map bandwidth into preconfigured circuits that are always the wrong size. RPT enables the provisioning of more services to more customers over the same cable backbone network. Transport bandwidth may be provisioned in exacting amounts with flexible burst capability.

RPT technology transports video, data, and voice payloads using MPLS [multi-protocol label switching] techniques. The video payload consists of MPEG data originating from either a QAM IF or DVB-ASI interface. The data payload is IP and originates from 10/100 BT, 100 BFX, GigE, or OC-n POS interfaces. The voice payload consists of T-1 or E-1 frames. RPT supports TDM voice by establishing network synchronization to an external Stratum clock (BITS, T-1). This allows transport of toll quality voice services originating from TDM equipment.

In conventional router networks, bottlenecks can occur when L-3 packet traffic exceeds the capacity of the router to efficiently "open" each packet, read the destination address, etc., and forward to the

next router. Today's solution requires the cable operator to upgrade to a higher capacity router.



RPT eliminates the layer three bottleneck and expense by operating in Layer 2-1/2 and using the RPT header for pack destination information. As each packet traverses network, the RPT header determines if the packet should be dropped or passed. RPT enabled boxes will likely reduce network dependence on high capacity routers.

RPT should better position the cable operator or its CLEC partner to sell differentiated services to end users [often B2B customers] through Service Level Agreements [SLAs]. A SLA is a contract between the operator and the user/subscriber setting expectations and pricing for transport services while establishing penalties for non-compliance. These end-to-end service guarantees include latency, bandwidth, and packet loss. For a successful business, SLAs require all traffic to be trackable, verifiable, and billable.

RPT enables this high-end service capability through the implementation of QoS and network management. Key to RPT's implementation of QoS is the preclassification of each [voice, video, data] packet as it ingresses the ring. Classification is based on "diffserve" code, source/

destination address, port, MPLS tag, or application or a combination of the above. RPT offers eight QoS classes from express forwarding [EF] for time sensitive services like video or toll voice, ranging to best effort [BE]. Once a preclassified packet enters the ring, software algorithms manage traffic congestion and properly balance add/drop/pass traffic at each node.

RPT's Network Management Control is standards compliant and includes the ability to control the provisioning of guaranteed services and implement SLA's. Point and click provisioning for bandwidth management, cross-connects, class of service, and accounting management facilitate the cable operators ability to

deliver, monitor, verify, and bill for differentiated IP services.

Summary

It is forecasted the growth and pattern of traffic in the cable operator's backbone network will migrate towards IP data. However, the MSO must also reliably and efficiently transport analog/digital video, and TDM voice services over this same network. RPT technology offers a packet level alternative that is efficient for all types of transport traffic, and allows the true economic benefits of a single platform transport solution to be realized.

IS THE DOCSIS CMTS SUFFICIENT FOR PACKETCABLE?

Burcak Beser *
Pacific Broadband Communications

Abstract

The DOCSIS 1.1 CMTS, as defined by the CableLabs, is almost ready for PacketCable. This paper discusses the additional functionality that will be required by a CMTS in order to satisfy the requirements of the PacketCable specifications. Equally important, this paper also identifies and discusses the external back office elements that are required for a successful PacketCable deployment.

This paper identifies the DOCSIS features that are used by PacketCable. Additional needs of telephony applications are identified, as well as an introduction to how these additional requirements are linked into DOCSIS features.

INTRODUCTION

PacketCable is a project instigated by Cable Television Laboratories, Inc. and its member companies. The PacketCable project is aimed at defining interface specifications that can be used to develop interoperable equipment capable of providing packet-based voice, video and other high-speed multimedia services over hybrid fiber coax (HFC) cable systems utilizing the DOCSIS 1.1 protocol. PacketCable defines a network superstructure that overlays the two-way data-ready broadband cable DOCSIS 1.1 access network. The initial phases of PacketCable cover only voice communications. This paper only addresses the issues regarding PacketCable

Cable Modem Termination System (CMTS) support for carrier grade voice over IP (VoIP).

PacketCable assumes operation over DOCSIS 1.1, which adds features to the basic DOCSIS 1.0 release capabilities in the areas of managing and packaging of the data services. PacketCable augments the basic back-office elements such as Data Provisioning and Data Management servers of DOCSIS, with comparable equivalents for VoIP services. A Record Keeping Server, Call Management Server, and Gate Controller are shown in Figure 1.

PacketCable requires the following additional protocols and functionality from a DOCSIS capable CMTS:

- Theft of Service Prevention
- Quality of Service
- Legislative Support

THEFT OF SERVICE PREVENTION

The basic assumption regarding theft of service prevention is that the Multimedia Terminal Adapter (MTA) is not resistant to customer tampering, and that the incentive to illegally obtain free service will lead to some very sophisticated attempts to thwart any network controls placed on the MTA. This customer tampering includes, but is not limited to: opening the box and replacing ROMs; replacing integrated circuit chips;

* e-mail: burcak@pbc.com

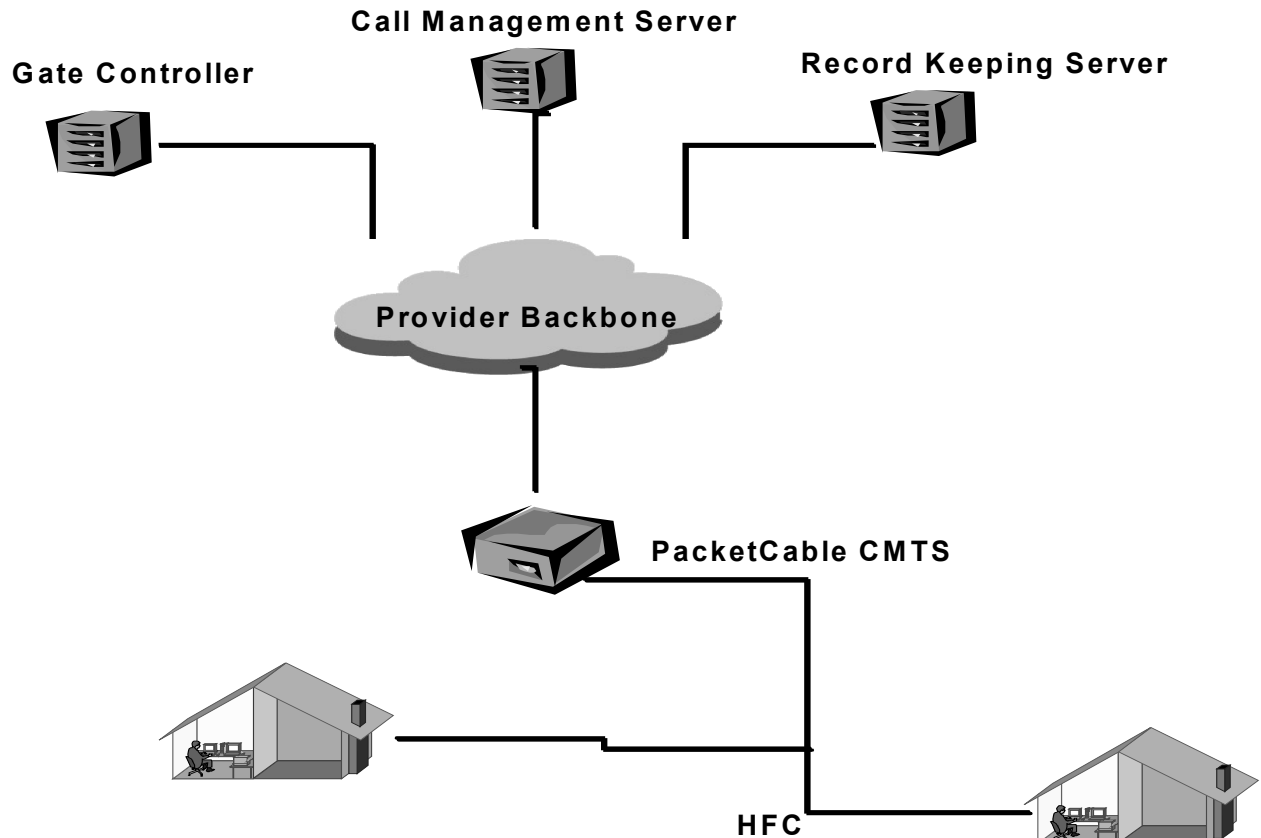


Figure 1 PacketCable CMTS and back office elements that are connected

probing and reverse engineering of the MTA design; and even total replacement of the MTA with a special black-market version. Examples of this degree of effort can be found in various other industries and technologies.

Since an individual MTA can be distinguished only by its communication over the RF network, it is possible, and quite likely that PC software may be written that will emulate the behavior of any MTA. In such a case the PC may be indistinguishable from a real MTA. In this case the software is under the total control of a customer.

Customers establishing high level QoS connections themselves

The MTA with sufficient intelligence can remember past destinations dialed as well as the destination address, or use some other mechanism to determine the IP address of the destination. It can then signal that destination itself (with some cooperation of the far-end client), and negotiate a high level quality-of-service connection via the RSVP mechanism or via the interface for an embedded client. Since no network agent is used in initiating the session, no billing record will be produced.

Even though the above scenario requires the cooperation of two altered MTAs, it is possible to achieve the same effect by manipulating/modifying only the originating MTA. If the originating MTA used the network agent to establish the session, thereby informing the destination in the standard manner of an incoming session, but again negotiated the high quality-of-service itself, there would be no billing record generated and the originator could obtain a free session.

Prevention of this scenario is accomplished by requiring per call authorization at the CMTS; without the proper authorization, any attempt to obtain carrier-grade quality-of-service to make the phone call will fail.

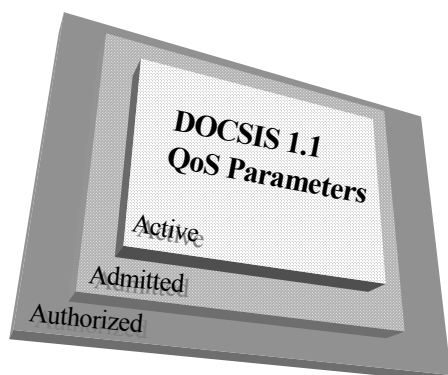


Figure 2 DOCSIS 1.1 QoS Envelopes

The DOCSIS 1.1 specification defines 3 envelopes that can be defined for an IP (service) flow: Authorized, Admitted, and Active. The “Authorized” envelope defines the boundaries of the QoS parameters that can be used by the initiator for specifying an IP flow. The “Admitted” envelope shows that a particular IP flow is admitted, and “Active” is the state in which the IP flow can be used with the QoS that was admitted. The relationship of these envelopes are shown in Figure 2.

This per-call authorization model requires the PacketCable CMTS to have an external means of setting the “Authorized” envelope by the Call Management Server (CMS). Since the Call Management Server knows when a call is in progress, the above-mentioned theft of service methods would be prevented.

The PacketCable DQoS specification includes the external interface that defines the “Authorized” envelope as Common Open Policy Service + (COPS+). The COPS protocol defined by the PacketCable has sufficient differences that a standard COPS client/server is not sufficient for the PacketCable needs.

The COPS+ protocol specification requires an extension to the objects and security mechanisms. The object that is manipulated by the gate controller via the COPS+ protocol is called a *PacketCable Gate*. The *PacketCable Gate* is a special construct that controls the “Authorized” envelope, and is at the heart of theft of service prevention in a PacketCable CMTS.

Customer alteration of the destination address of voice packets

Another theft of service scenario is shown in Figure 3. Two remote MTAs that are far apart, each make a local call. Once the bandwidth and connection for these local calls has been established, the MTAs change the destination addresses to cause their VoIP streams to point to each other. The billing system continues to bill each of them for a local call, while the customers are actually engaged in a long distance call.

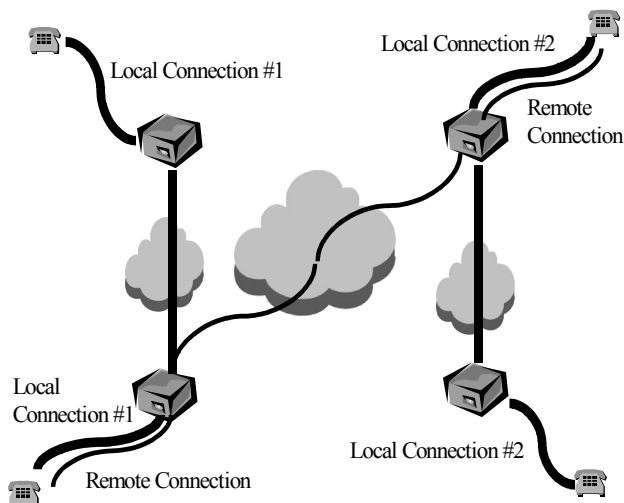


Figure 3 Using two local connections to achieve a remote connection

The solution to this theft of service mode has two parts. First, the previously defined authorization interface (PacketCable COPS) now should pass information regarding the IP addresses of the far end. Second, the CMTS should police the IP flow for compliance.

The DOCSIS specification does not mandate the use of egress policing of classifiers for the DOCSIS CMTS. A PacketCable-capable CMTS should implement the optional feature of policing egress IP flows.

MTA non-cooperation for billing

One can easily imagine what would happen if there was a message from the MTA on session establishment that said, "OK, called party has answered, start billing me now," or a message on hang-up that said, "Session has completed, stop billing now." However, there are more subtle ways that a user could have the same effect as tinkering with such messages if they existed.

With the current PSTN scheme, users are billed for the entire timeframe that they spend actually connected, but they are not billed, for example, for the 30 seconds that the far-end phone was ringing. As shown in Figure 4, the billing must be connected to actual VoIP QoS usage. Whenever the PacketCable CMTS activates a service flow, which is connected through a *PacketCable Gate*, it sends an event message to the Record Keeping Server, which is later used for billing purposes.

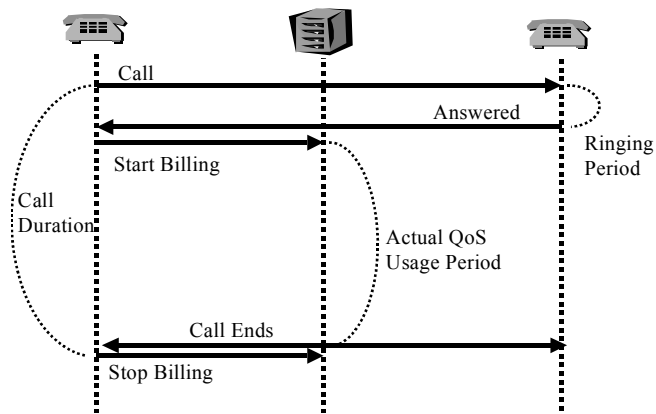


Figure 4 Call Billing

This state by which the connection is made and billing started is referred to as *PacketCable Gate Open*. The same event messaging occurs when the created service flow is deleted at the end of the call. This behavior is different from the DOCSIS standard where no messages are generated on the WAN interface when a service flow is activated or deleted.

The PacketCable specification defines the above-mentioned event-messaging protocol as Radius. The Radius protocol defined by PacketCable has sufficient differences from that of a standard Radius client/server and is not sufficient for PacketCable needs, and an additional content and security mechanism is specified by PacketCable.

Use of half-connections

In this theft of service scenario, one of the MTAs (originator) in the call signals the start of conversation but the MTA at other end (termination) does not signal the start of conversation at its end.

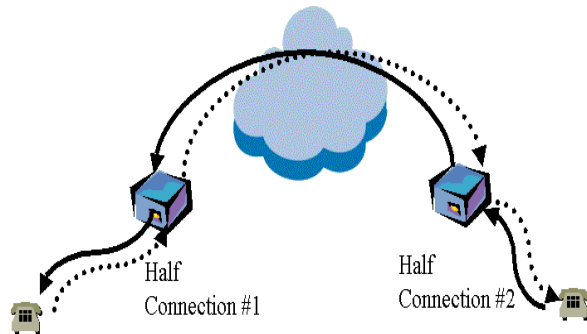


Figure 5 Using two half connections for bi-directional call

In this case, only one *PacketCable Gate* is opened, and the users and network are left with a “half-connection”. Given that the terminating MTA did not send a “conversation-started” message, the network cannot legitimately bill the user for the half-connection. However, it is possible for two colluding clients to set up two half-connections, neither of which is billable, which can be combined to give a full connection between the two parties.

This is an example of theft of service that could occur in the absence of resource use coordination at both ends of the call. Coordinating the operation of the two *PacketCable Gates* can prevent fraud of this type.

Whenever a *PacketCable Gate* is opened (thereby granting carrier grade QoS on the Cable segment) a message is sent to the far end *PacketCable CMTS* that the half connection is started, and if it does not receive a message that indicates the other end has committed to the call by some certain time then the call will be terminated.

The *PacketCable* specification defines this *gate-coordination* protocol as Radius with *PacketCable* specific content and security mechanisms.

The theft of service attack defined above can also be executed by early termination of the connection by one end only. In this case the Records would show very short call duration. This mode of attack can easily be prevented by coordinating flow deletion.

PacketCable defines that the same *gate-coordination* protocol used above will be used for flow deletion as well.

Fraud directed against unwanted callers

Due to details of the call setup sequence, it is possible that bandwidth authorization at the destination will be more generous than that at the source. Given this, it is then possible for a called party to reserve and allocate bandwidth far in excess of the final negotiated amount, resulting in the calling party being charged more than expected. If available, this technique would likely be used against telemarketers by individuals in response to unwanted calls during dinner.

For this reason the *gate coordination* defined by *PacketCable* includes the examination of the bandwidth that is being used.

QUALITY OF SERVICE

The use of DOCSIS 1.1 enables PacketCable systems to provide carrier grade QoS for VoIP communications. Some of the properties of DOCSIS 1.1 that are applicable to MTAs can be summarized as:

- Multiple service flows, each with its own class of upstream traffic
- Both single and multiple voice connections per DOCSIS service flow
- Prioritized classification of traffic streams to service flows.

For a DOCSIS 1.1 CMTS the following are the QoS settings that play a crucial role in providing carrier grade QoS:

- Guaranteed minimum/constant bit rate scheduling services
- Constant bit rate scheduling with traffic activity detection service (slow down, speed up, stop, and restart scheduling)
- DOCSIS packet header suppression for increased call density
- DOCSIS classification of voice flows to service flows
- TOS packet marking at the network layer
- Guarantees on latency and jitter
- Reclamation of QoS resources for dead/stale sessions
- Two stage QoS setup

The primary mechanism for providing low-latency quality of service for media streams in the access network is the DOCSIS 1.1 flow classification service. This service classifies packets into specific flows based upon packet fields such as IP source and

destination addresses and UDP port number parameters. In the upstream, such classified packets are transported via an appropriate constant bit rate service (for current codecs), as dynamically scheduled by the CMTS. In the downstream the packets are transported via an appropriate high priority queuing and scheduling mechanism. DQoS (between CMS and CMTS) and DOCSIS (between CMTS and CM) signaling mechanisms are used to dynamically set up the media stream flow classification rules and service flow QoS traffic parameters.

Providing Timely Call Setup

During call setup a number of messages are exchanged between various entities. Of these messages, the QoS setup messages are handled by the CMTS. Since the QoS messaging takes place in real time while callers wait for services to be activated, the protocol that is used for QoS setup must not impose unnecessary delays. The number of messages, which traverse end-to-end, should be minimized. For this reason PacketCable partitions the resource management into two distinct segments: access and backbone. Segmented resource assignment is beneficial for two reasons:

- It allows for different bandwidth provisioning and signaling mechanisms for the originating network, the far end network, and the backbone network.
- It allows for resource-poor segments to maintain per-flow reservations and to carefully manage resource usage. At the same time, when backbone segments have sufficient resources to manage resources more coarsely, it allows the backbone to avoid keeping per-flow state, and thus enhances scalability.

DOCSIS only specifies QoS control for the cable segment through DOCSIS MAC messages. DOCSIS MAC messaging is useful only if the QoS requesting entity is directly connected to the RF link. If a client is connected through an IP network behind a cable modem, then that device cannot request QoS from a DOCSIS CMTS.

For this reason, the PacketCable specification uses RSVP+, which is a modified version of the RSVP protocol. Using this modified RSVP protocol it is possible for all client devices to request high quality IP links from the PacketCable CMTS.

For the IP backbone, the PacketCable specification is more relaxed and allows the use of a number of protocols including RSVP, DiffServ and Aggregated RSVP. Due to limitations on the scalability of both RSVP and of DiffServ, it is expected that aggregated RSVP will be used for signaling IP backbone high quality links.

Providing low delay voice transport on a Cable Network

For voice services, the end-to-end packet delay needs to be small enough that it does not interfere with normal person-to-person interactions. For normal telephony services using the PSTN, the ITU recommends no greater than 300 ms roundtrip delay. Given that the end-to-end backbone propagation delay may comprise a significant percentage of this delay budget, it is important to control delay on the access channel, at least for long-distance calls.

DOCSIS 1.1 specifies a new scheduling service called Unsolicited Grant Service (UGS) in order to reduce the delay that is

introduced by the upstream cable segment. Unfortunately just the existence of UGS is not sufficient to minimize upstream cable network delay. The generation of the VoIP packets has to be synchronized to the UGS schedule, which is itself synchronized to the DOCSIS timestamps.

On the far end the VoIP client will play out the incoming data stream. The far end can be another VoIP client or can be a PSTN gateway, which converts the VoIP packets into a PSTN data stream (a DS-0).

In PSTN networks the voice samples are transmitted using common reference timing to which all the PSTN gateways should be synchronized. Due to the sampling synchronization, the incoming frames must be synchronized to UTC as well. If the synchronization does not take place, the playout buffer would underflow or overflow over time, which would result in pops and clicks and would effect fax/modem communications carried as VoIP data.

The DOCSIS specification does not mandate that the DOCSIS CMTS be synchronized to any external clock source, and due to frame synchronization issues, the PacketCable CMTS should be synchronized to the PSTN common timing reference.

LEGISLATIVE

The PacketCable system should be ready to adopt the same legislative requirements that are in place for PSTN systems. Two of these requirements, electronic surveillance and privacy, are specific to a PacketCable CMTS.

Even though it is possible to achieve compliance outside of the CMTS, the PacketCable CMTS is at the right place to achieve compliance without unnecessary overhead.

Electronic Surveillance

Electronic surveillance includes both interception of communications and the acquisition of call-identifying information. Since the call-identifying information is not retained in the CMTS, only communication interception should be handled by the PacketCable CMTS.

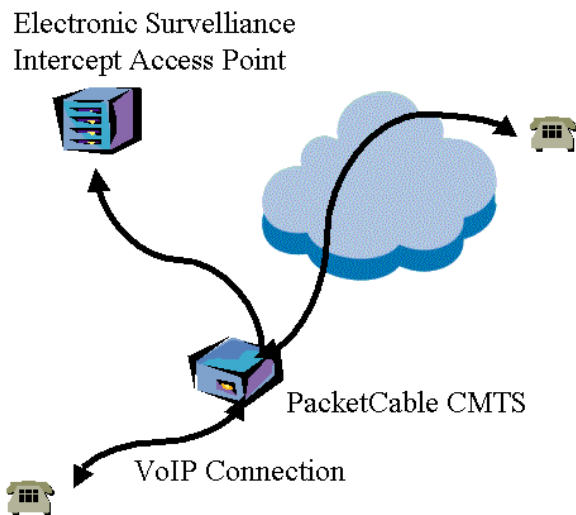


Figure 6 Electronic Surveillance using PacketCable CMTS

A PacketCable CMTS that supports electronic surveillance is responsible for duplicating the VoIP packets of a specific IP flow and sending them as defined in the PacketCable Electronic Surveillance Specification to a predetermined end-point, such as to a local Electronic Surveillance Intercept Point at a local federal law enforcement office as shown in Figure 6.

Privacy

The privacy issue on VoIP systems has multiple dimensions: One dimension is the telephony representation of caller ID, which is achieved by careful design and implementation of the PacketCable call signaling protocol.

The second dimension is IP address privacy. This issue arises from the fact that the IP address that is contained within IP packets can be used to determine the location of the caller. Today there are multiple systems that can accurately pinpoint the location of any given IP address.

IP address privacy can be achieved if the PacketCable CMTS performs Network Address Translation on VoIP packets.

There is a beneficial side effect that comes from double NAT: The telephony devices do not have to be globally routable. Having telephony devices that are in the private address space would help greatly to alleviate the depletion of routable IP addresses, especially if one thinks that there would be millions of such devices.

IN SUMMARY

Even though DOCSIS 1.1 is an essential baseline for PacketCable support, a DOCSIS 1.1 CMTS alone is not sufficient for supporting carrier grade voice transport over the Internet.

The PacketCable CMTS needs to support additional protocols and functionalities for the areas of theft of service prevention, carrier grade quality and legislative issues.

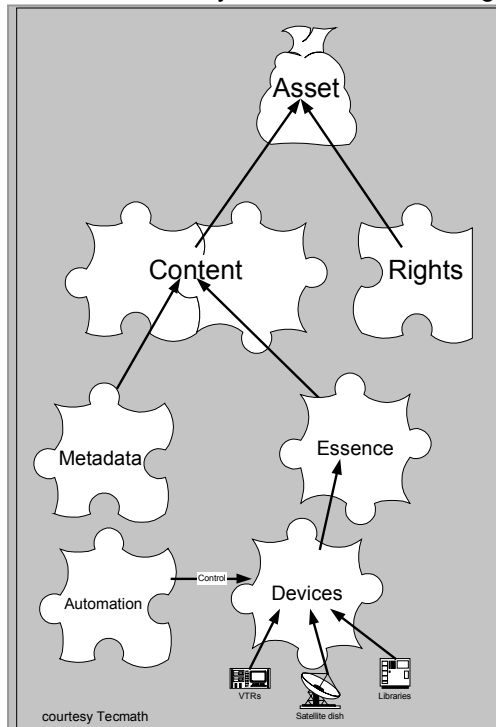
MEDIA ASSET MANAGEMENT: THE NEED, THE CHOICES, THE PAY-OFF

Wendell Bailey
Chief Technologist
Advanced Broadband Technology
NBC

Dr. Peter Smith
Vice President
Technical Planning & Engineering
NBC

Introduction

NBC, as a media production and distribution company has for some time been concerned about the need to make its various media assets secure as well as easily and economically available to those in its operating groups who have a regular need to access it. In this paper, It is our intention is to lay out most of the operational and technical issues that need to be considered by any media asset owner. The MAM system described is a digital



content management system for use in broadcast production and archiving. The details are the result of a 12-month study of NBC's needs and capabilities.

Terminology

Most media companies currently have systems in place to handle automation, devices, essence, metadata, content, rights and assets. Unfortunately, there is scant coordination between these elements today, which are mixtures of analog and digital systems. The question before us is; how do we create a unified digital asset management system that provides timely access to assets, increases the value of assets and creates efficiencies in the production and distribution of programming, to the degree required for investment in the system?

An ASSET is defined as CONTENT for which one has the RIGHTS to sell or distribute.

CONTENT is defined as ESSENCE and METADATA used to describe that particular ESSENCE.

ESSENCE is the raw video, audio, or image from devices like video recorders, video servers and satellites.

History of the Project

In December 1999, NBC began a company wide effort to study Media Asset Management

(MAM) systems, capabilities and requirements. The MAM Group's mission was to determine if there was a need for better asset management tools within NBC and if so, what functions would they perform and what benefits would they offer.

Baseline Functional Needs

The results revealed pockets of Media Management expertise, a mix of different media asset management systems and a distinct desire for greater Media Asset Management capabilities. Throughout production divisions, users chiefly articulated desires to:

- Improve Workflow Productivity by having media more readily available
- Communicate, Collaborate and Gain Access to Assets across & within businesses
- Exploit New Business Opportunities by marketing latent assets

NBC's Media Asset Management Model

The MAM group developed a stratified architecture for a media asset management system consisting of 3 layers:

- An End User Application Layer – reflecting current operations in a broadcast environment like shooting, editing, integrating with commercials and playing to air.
 - A Foundational Layer – reflecting the need for standards in digital formats, networking and storage, so that baseline functional needs could be met.
 - A Middleware connection layer – recognizing that equipment and formats must be flexible in the business units to meet particular needs, but must link to a common digital format suite in the media asset management foundation.
-
- Scalable, Extensible, Redundant MAM system
 - Metadata Structure must be Defined and Standardized
 - Visual media management references must be Accessible From the Desktop
 - Complex Technology should be Invisible To User

- Flexibility – The system must accommodate existing and future standards and formats
- A Digitization Strategy for New & Existing Assets is required
- Fast Production Cycle Times – nothing should impede current production speed
- Copyright Protection
- Simple Metadata Creation Processes
- Open Architecture
- Security Technology needs to be incorporated

NBC News Archives

In 1996-97, NBC News Archives created a database and search system, OMAR, for the 2.3 million tape/film inventory in the NBC news archive. OMAR (Oracle Media Archive Research) is a networked, text based research tool used by nearly every business unit in NBC.

Out of the hundreds of hits an OMAR search may produce, there is little information that allows an archivist to rapidly winnow out the unsuitable and order the most highly desired clips. Hundreds of tapes might need to be pulled from the library, often dubbed and delivered to the producer who must go through each to find a suitable clip. OMAR provides valuable handles to the contents of the archive, but the archive is so rich, with multiple clips on each tape that a lot of time is spent searching and very little time finding.

In addition to the lack of visual information in search results, News Archive was interested in capturing production data developed by producers, writers, editors and others in the production process. This information would be valuable to archivists in appropriately describing captured clips, so searches produced more relevant results.

End User Applications

In the Media Asset Model, the application layer should augment current production methods with digital tracking and management information. It should allow each business unit to select specific equipment and components that help them meet their objectives and provide additional functions of sharing material, creating new products and streamlining the workflow.

- Business objectives serve to dictate equipment and methods for production. For example, Nightly News shows have very short turnaround times, measured in hours and minutes. News shows receive much of their production material from feeds received by satellite from remote bureaus or locations. Local news production (at least at the 13 NBC stations) works on similar timelines but have news crews bringing tapes back to the station throughout the broadcast day. The media asset management system needs to create efficiencies but respect diversity in program production and efficiencies already in place.

Acquisition

Acquisition of programming has two meanings. It can be the commissioning of programming, such as entertainment shows or topical magazine shows contracted from an outside production company. It is also the highly decentralized activity of capturing video, audio and data in the field.

Commissioning & Pre-production: Storylines, storyboards, locations, crews, scripts are assets. They can be managed by a properly designed system. As a metadata standard is defined, and digital equipment is deployed (PDA's, wireless modems, disk camcorders) the acquisition stage can be integrated into the overall asset management system.

Capturing: The wide variety of crew responsibilities and methods used in this activity preclude short-term changes. We expect to see digital recording in the field, on removable optical disks, possibly by 2003. The adoption of digital acquisition equipment will take place over a 10-year period, from 2003 to 2013. The digital asset management system we design today must easily accommodate the integration of digital acquisition systems. Acquisition equipment vendors need direction to add content management features to their product and the publication of metadata standards to these vendors may help to provide needed direction.

Facilities recording feeds on video servers are acquiring digital essence. The MAM system will add metadata to the essence, during ingest, creating content.

Ingest

Ingest is the digitization of analog materials. Video, audio, paper notes, contracts, time and date information is analog until ingest. Ingesting activity is currently ad-hoc. A media asset management system places a structure on this activity and must be carefully engineered to produce efficiencies. Ingest is the first opportunity to add metadata to essence that is ingest. The MAM system described herein pre-supposes that metadata will be added at this ingest stage. A specific question for each business is whether to ingest high resolution, broadcast quality video and low-resolution proxies, or simply low resolution and metadata. This decision is budgetary, as high resolution files can be expensive to store on hard disks. The chosen architecture needs to accommodate both digital high-resolution files and analog videotape footage, so that unreasonable economic demands are not placed on business units.

Ingest changes the basic workflow of production units. Resources currently dedicated to backtracking to gain information on essence can be deployed up front, at ingest, to support the desired functions of sharing material, creating new products and streamlining the workflow. Ingest can also include the current activity of Producers and associate producers, who screen material and make logs of tape contents.

Ingest is well understood by archivists, who add metadata to essence they catalog. Their critical to quality needs are listed in bullets below.

- Graphical User Interface must emphasize editorial considerations, not on navigating around the computer application.
- The Synopsis, Subjects, Personalities, Places, and Dates fields should be displayed proximate and readily accessible.
- The process of associating the content descriptions [metadata] to the assets should be fluid.
- The entry and validation of descriptors should be friendly, perhaps like a spell-check operation.
- There must be a Standard English spell check.

- Considering the possibility that a server, network, or workstation could crash, users must be able to save content and order data without closing their records, and the system itself must automatically periodically store unsaved data.
- It must be possible to create a new distinct record, which carries the Editorial ID and selected other information from the previous record, forward to the next record.
- It must be possible to generate ID numbers based on specific editorial information. (Our ID numbers, for instance, are not crystallized at random, but rather are synthesized by combining bureau, assignment, and date information with a sequence number.)
- It must be possible for users to pre-set certain search criteria.
- Rights Data Linked to Other Footage Records. Thus, the same set of rights does not have to be re-entered for each instance.
- It should be possible to scan hard documents, e.g., third-party contracts, into the database.
- Users with data-entry privileges must be able to instantly propose new descriptors that are immediately, if tentatively, associated with that record.
- Set Resolution of Thumbnails.
- Ability to Set Thumbnail (snapshot of video essence) frequency at Ingest.

Production

Tape logging, scripting, tape dubbing, ingest into nonlinear edit systems, editing, graphics creation, effects creation, screening are subdivisions of production activities. A MAM system will affect production activities in different ways so we describe them separately.

Tape logging

Currently, operators put a tape in a screening VTR and play it, writing time code numbers and indicating shot changes, noting the contents of a videotape recording. They often capture data on a piece of paper, which is then handed off to an editor and sometimes stored with the tape cassette.

With a MAM system, tape logging becomes part of an ingest function. Operators will be able to concentrate on simply adding higher

level information to data captured automatically from the video and audio contents of the tape. During the ingest stage, low-resolution images will be captured and linked to the high-resolution video that spawned them. The logging process will produce a universal ID to the clip which can be referenced back to the high resolution videotape by reel number barcode and timecode on the tape. All videotapes and datatapes used in a MAM system will need a consistent barcode numbering scheme. As producers screen the low-resolution files, they can continue to add information about the essence that will help others to later locate clips they may have an interest in using.

News Room Computer System Scripting

It is common for writers to use an I-news system (AKA - Avstar or Basys) to write a story using interview transcriptions and research information. They guesstimate timings and expect the editor to find enough video filler to support the story.

With a MAM system, writers will be able to search for and recall ingested low-resolution video in a window on their word processor. They can look at the video they are writing to, decide on appropriate wording and read time, and roughly piece together interview parts that convey the meaning of their story. MAM is expected to allow the production of more relevant text, and improve the speed of production by reducing the amount of re-writing and re-editing in the current process.

The newsroom computer system also contains valuable metadata that can help archivists.

- **Producer Data Entry.** A mechanism by which production staff can enter their screening notes directly into the Archives database (for later acceptance or rejection by Archives staff) would be beneficial.

Tape dubbing

Currently, multiple copies of studio feed tapes are simultaneously recorded. Camera field tapes are one of a kind and if multiple producers want to work on related stories, dubs of field tapes must be generated. Also, if archive tapes are one of a kind, they may be

duplicated before leaving the archive. In summary, a lot of time is spent making dubs and moving tapes to satisfy production needs. This not only slows down production while dubs are made, but adds transportation dependencies. Using a MAM system, the ingest stage could create low-resolution files which could be accessed within a few seconds of recording. Several users could then access the material without duplication. The material is screened over a digital computer network, minimizing the need for transportation logistics. If many users (30 or so) must screen material simultaneously, very low resolution duplicates of the file can be automatically generated and put on an alternate server and IP network. The need to dub is nearly eliminated by using a properly designed MAM system.

Ingest to nonlinear edit devices

Differences in Nonlinear edit (NLE) systems impacts the efficiency of overall digital systems. For the most part the systems are used as black boxes, which require videotape, input and output. At the same time, the NLEs create digital media, provide ample amounts of metadata useful to archivists and can control videotape recorders for their ingest. A MAM system will need to “glue” these legacy systems into its architecture. Ways in which this gluing can be accomplished include:

- Ingesting high-resolution file formats that are compatible with the NLE systems, so that the ingested files can be instantly transferred to the NLE storage.
- Ingesting low-resolution formats that are compatible with the NLE systems, so that the browse files can be transferred to the NLE storage at faster than realtime.
- Capturing metadata that can be transferred to the NLE for use by editors and the NLE system for batch input
- Creating edit decision lists with producer's notes in desktop browse editors that can be ingested and used by the NLE system
- Capturing low-resolution video with high-resolution audio for direct ingest by the NLE so that edited audio tracks do not have to be re-digitized and remixed.

- Using the MAM Robotic Library to load video tapes required by the NLE's batch input function
- Wiring the Video, audio and timecode from a MAM Robotic Library's Videotape machines to the ingest system of the NLE
- Allowing the MAM database to query the NLE's database to gather information about
 - clips that will be archived
 - clips that will be used on air
 - clips that have restricted rights
- Putting data tape drives on NLE storage systems that are compatible with MAM's Robotic Library system

Gluing legacy NLE black box silos into a standard architecture presents one of the most difficult technical challenges to the MAM design. The MAM multigenerational plan should include these requirements to insure the smooth integration of high resolution, highly functional nonlinear tools into its architecture.

Editing

Audio layering and mixing, graphics creation, video effects creation, frame accurate timing, video level and color correction, and technical quality control are core aspects of the editor's craft. It is unrealistic to expect producers to take responsibility for these techniques. Pre-editing functions like shot selection, shot timing, shot ordering and archive retrieve are primary targets for the efficiencies of a digital MAM system. A MAM compatible edit system should therefore be as tightly coupled with ingest, logging and scripting functions as possible, to eliminate duplication of efforts.

Creating graphics

Complex layered graphics are often created for show opens, teases, transitions and promos. The collaboration of the edit teams with the graphics team results in the need for file exchange between the devices. As another example of the inefficiencies of legacy digital systems, there is currently little or no capability for edit systems to ingest a high end graphics file directly. While some efforts are underway

to build compatibility, the MAM system should probably look, short term, to legacy digital video routing systems and intercoms to link graphics with edit systems.

Interoperability of graphics systems files from vendors like Quantel, Chyron, MacIntosh, and PC software vendors should be easier to accomplish.

Creating effects

Often accomplished in NLEs or video switchers, effects are generally transitional, but also create "new clips". Clips that specifically relate to rights management, such as those that have had a face obscured or a voice disguised need to be tracked by the MAM system. It will be necessary for operators to properly note new clips and make the system aware of these derived clips to properly track rights and usage in the enterprise.

Screening and approvals

Screening is usually done for review and approval on a TV monitor in standard resolution from a videotape dub of an edited piece. The number of approvers who must screen and pass a piece for distribution varies widely throughout production units. Local news may require no screening approval beyond the producer-editor's. A magazine type news show investigative piece may require screening up to the CEO level, involving 30 or more people over 2 or more screen processes. Advertising and promotions screenings are similarly complex, with bi-coastal production and often, the approval of outside producers and agents required for distribution. Security issues complicate this rigorous process, and pre-approved, pre-air pieces must be carefully handled.

A properly designed MAM system will allow secure screenings of digital video files. While it does not have to be included in a baseline MAM system, specific digital networking software is available to deal with approver's notes and comments, embedded in the digital file itself. With such a system, revisions can be more accurately accomplished. The MAM network and security will be robust enough to allow screeners to see and hear the piece they need to pass or fail. Each business unit will be required to assess their need for a more robust

screening system and integrate that system into the MAM system, if it is required. In any case, MAM needs to track the approvals, whether done in an analog screening room or on the digital workstation of an approver. The MAM metadata structure will allow users to query the approval status of a project, for accurate assessments of a piece's status.

Archive/Retrieve

The MAM system should be designed with both on-line and off-line digital storage to reduce the cost of storing large essence files and backing up smaller metadata files. Archive and retrieve refers to the process of putting digital files onto low cost removable media and getting files off the media and on a hard disk again.

There are many types of files in a MAM system, including low-resolution browse, metadata, database, graphics, scanned documents, high-resolution video and audio files, and very low resolution video files. Different file types may require different types of removable media in a robotic media handler. It is important that the system be able to handle high-speed media like tape drives, and quick seek media like optical disks. The system should be compatible with videotape drives as well as datatape drives. As the archive system is the main repository of digital data, rights management must be embedded in the retrieval process:

- Rights information must be readily visible, and the very fact of an associated restriction should be prominently displayed along with the record.
- Non-Archives staff should be able to place orders directly through the application, which then must be approved by archivists. The Approval screen will present the Archives all essential information upon which to make a decision, e.g., duration, rights, original format, etc.
- Any Restrictions should appear during the ordering process, as a warning. [A report might also be generated which alerts Rights & Clearances of any requests for Restricted materials.]

- Ease of use for untrained users:
 - Untrained User's Search should be Straight-Forward The screen used by untrained users (which could be the same one used by trained users) must be inviting and must be able to parse and execute a simple query using keywords and dates. The more flexibility we can give untrained users *without* adding confusion to the screen, the better.
 - All Text Fields Query. There must be a query that searches across multiple select fields as if the contents of those fields were all in a single field.
 - Weighted Results. It would be helpful if the novice researcher did not have to think about Boolean operators but could simply enter terms and the search engine would then present results under ranked classes, e.g., records which contain the search term as a phrase, records which contain all the terms but not in a phrase, and records which contain some of the terms. A simple numerical ranking system (such as is used on Internet search engines) would not be adequate.
- Flexibility of Search Interface for Trained Users.
 - Users should be able to apply Boolean logic to virtually any field, including fields on related tables. An advanced user should have virtually total flexibility in how they structure their query.
 - Sorting. It must be possible to sort the hit list (regardless of scoring) by event date.
 - This Day in History. It should be possible to search the date field for month and day leaving the year a wildcard.
 - Searching by Paragraph. The ability to query by the intersection of search terms *within the same sentence or paragraph* rather than merely within the same record is an added plus.
- Computer-Enhanced Searches— e.g., those allowing natural language queries or dependent upon thesauri hierarchies — must return results which reflect the quest (a query for *secret agents and moles*, for instance, should not return stories about federal wildlife officers and little animals) and should not miss appropriate records.
- Useful Presentation of Hits. All pertinent information should appear on the same screen, including:
 - Physical Material Characteristics and Locations. All physical material associated to an editorial record should be readily seen. If material exists on film, we should easily see if a film transfer already exists. Also, all film attributes should be apparent (type of film, length, b&w or color, silent or SOF). If there is no physical material associated to the editorial record, we should see that immediately. And the provenance of each dub or transfer should be immediately obvious through this screen.
 - Checkboxes which allow user to choose inventory without navigating.
 - Availability and, if on loan, then to whom, etc.
 - Search Terms should be highlighted in the display.
 - There should be a "Find on this Page" and "Find Next" search feature.
 - There should be several different views of the hit list one of which shows just a few fields of each record, one of which shows the full record, one which shows thumbnails, etc.; ideally, the user could select the fields which would appear in the view, and displays the search term in context — i.e., shows the ten words before and ten words after the search term.
 - Printing:
 - Full Record or Hit List. It should be possible to print either the full record (all fields) or just a hit list or a user-determined set of fields.
 - Highlighting. The search terms should be highlighted in the printout.
 - E-Mail. It should be possible to E-Mail or Download the results.
- Grouping the Selections. It should be possible to sequence and review selected materials, and possibly to sort into distinct bins for different projects.
- Ease of Marking Selections and Placing Order:
 - Marking Ins and Outs. It should be possible to select Ins and Outs from

- streamed video and audio, and to calculate the total running time.
- Display resolution of thumbnails can be set Lower than Ingest Resolution. [4]
- Inventory Management
 - Tracking Asset Movement. (It must be possible to precisely track tape and film locations – to check tapes out to borrowers and check them in when they're returned, to note when tapes are on a shipment and when the shipment has been received, and to check whole shipments out to a single person.)
 - Displaying Inventory Information. (It must be possible to display all individual tape information on a single screen, including current location, current loan information and history of loan and shipment information.)
 - Overdue Notification. (It must be possible to send e-mail notifications out periodically to borrowers that their loans have aged excessively.)
 - Deaccessioning. (It must be possible for authorized parties to remove footage from the collection but to maintain a record in the database of that action.)
 - Transferring Checkouts. (It should be possible to expeditiously transfer checkouts between borrowers.)
- General Administration
 - Maintain Accounts and Categories:
 - The application administrator should be able to establish user accounts, including setting numerous security/role levels, without the intervention of a database administrator.)
 - The application administrator should be able to set up Bureaus, Domains, Assignments, etc., in the system, without the intervention of a DBA.
 - Create Reports both ad hoc and regular). Application Administrators should be able to query any table(s) and present results.
 - Maintain Thesaurus, including the ability to merge heretofore distinct descriptors into one, to approve or remove descriptors, to modify them.

- Create Shell Records. Application administrators should be able to create series of shell records (setting a range of Editorial ID numbers).
- Make individual corrections. Users with sufficient privilege should be able to query the database, produce a result set, and modify individual records from those results without first navigating to a separate input screen.
- Make bulk corrections. Application Administrators should be able to make mass corrections across records.
- Maintain Rights and Clearances Data. Rights and Clearances staff should be able to maintain and update Rights and Clearances tables.
- Security Levels. Application Administrators should have maximum flexibility in setting various user privileges.
- Usage Statistics. Archives administrators should be able to know who's logged in when, for how long, amount of streamed data, etc.
- Messaging. System administrators should be able to alert active users on-screen when, e.g., the system is to be taken suddenly out of service.
- Remote Access. System administrators should be able to mirror a user's session in order to remotely assist users having difficulty using the system.

Integration

Integration is the joining of pre-produced elements into a "seamless" stream of content and commercials. Integration also refers to the customization of output products, for specific distribution channels. For example, MSNBC's show "Time and Again" is finished in different ways for distribution through MSNBC, NBC Network, and Airline viewing. Different "bugs", distributor ID's, opens, closes or credits may be required for each distribution channel. Integration backs into the production process in the creation of pre-produced elements and will only become more complex and time consuming as digital channels proliferate. Content owners and program producers as well as business leaders indicated their desire to create new products, and this stage is a primary location for enabling new products, and avoiding prohibitive additional costs.

A MAM system can aid integration by:

- Automatically timing show segments for commercial integration
- Routing show segment files to servers for further customization
- Generating customization templates that automation systems can use to create customized products
- Extracting vital show information for distribution to Electronic Program Guides
- Extracting graphical elements and audio used in the production chain to generate lower resolution files for internet and wireless distribution

Commercials

Commercials are probably the most important aspect of the integration phase, and indeed, the business itself. In the typical station environment, the Traffic Department handles commercials. Commercial time slots are sold with provisions and guarantees for audience size, as rated by Neilson and others. As media asset management systems provide new functions to interactive consumer systems, commercial sales and delivery systems much match new functions, to boost income streams from improved interfaces to our customers.

Control of the viewer's screen becomes all important, which can only be accomplished via a 1 to 1 relationship with the consumer, or, providing information to the distributor in real-time, that allows a distributor to insert commercial messages at appropriate times. Feedback from customer sites or tight control of the distribution timing is critical to the control of the screen.

Distribution

Distribution is staging product for delivery to the customers. As the distribution of media becomes more complex, with analog on-air, digital on-air, cable, internet, broadband, wireless, data etc., having handles on digital products, such as those provided by a MAM system are becoming more important.

Automation systems will be forced to incorporate more complex associations, including advertising links, banner integration and digital headers in digital products. Return information, like viewer preferences and viewing times will be used to optimize distribution strategies. A MAM system will help programmers and sales staff to analyze programming for best possible results.

Foundation Layer

Clearly there must be standards established for an enterprise wide archival system to be successful. The difficulty is choosing standards that are useful to all users of the enterprise, and are also robust and extensible.

Format of Digital Suites

Archive and Distribution format MPEG-2

MPEG-2 SP / ML provides adequate quality in standard definition for distribution and archive. It is an open decoding standard, which means that future generations will be able to read encoded files without regard to decoding IP. It is very efficient, high quality, component based and can be encoded in realtime. It is efficient because of temporal compression, which unfortunately means it cannot be edited in its native form. It can be edited if it is decoded and re-encoded, which degrades the signal quality and risks audio impairments. MPEG-2 has a well defined "transport stream" which allows satellite distribution of interleaved video, audio and data signals. The MPEG-2 transport stream can be ingested and used in limited integration applications by some video servers.

Production format - DV

DV based compression, like DVCPRO, DVCAM, and consumer DV are editable in their native forms, don't degrade when edited, don't degrade substantially when decoded and re-encoded. They are at least three times "heavier" than MPEG-2 at 8 Mbps, so, to archive in DV format is not as efficient as

archiving in MPEG-2 long GOP. (We don't edit in the archive).

Browse Format – MPEG-1

The browse/edit format included in the above proposal is MPEG-1 at about 1.5 Megabits per second. This is an open format used in Video CD and other desktop applications. If encoded properly, the quality is equal to S-VHS, which can be put "on-air" in a pinch. The 1.5 Megabits per second limits the switched 100-baseTX network capability to about 30 clients, using a single server simultaneously. Additional servers and networks are required to reliably expand MPEG-1 browse editing to more than 30 clients.

Alternative Browse Formats

MPEG-4

MPEG-4 will be a scalable standard, starting at about .3 Megabits per second and allowing data rates up to 4 Megabits per second. This will allow the single browse server, which currently supports 30 MPEG-1 users, to support 120 MPEG-4 users. The MPEG-4 signal can carry digital "objects" which could be graphics or metadata, so users could display a data overlay with running picture. MAM system upgrades to MPEG-4 may involve IP licenses for MPEG-4 encoding that increase upgrade costs, MPEG-4 formats have been engineered for lower bitrate channels, such as those involving the internet and low cost media like DVD. MPEG-4 could form a standard for internal production and external distribution, if jog audio and frame accurate timecode are properly integrated into MPEG-4 formats.

Wavelet

Some companies are also developing a scalable Wavelet compression browse format. The system is marketed as IPV for Internet Pro Video and is aimed at wide area production. Wide area production is not currently used in production circles because of bandwidth limitations and costs of Wide Area Networks. IPV is not a standard, open format, but holds great promise if face to face production meetings are not a requirement and production time is not measured in minutes. Using an IPV system would allow editors anywhere in the

world to quickly edit footage that is arriving on a server anywhere else in the world. Wavelet compression is attractive because of its scalability, but the lack of International Standards supporting the format currently limits its usefulness in broadcast applications.

Screening format - RealVideo

The basic client is free, the encoder and server are cheap. MPEG-1 browse files can be used to encode RealVideo in faster than realtime. In large MAM systems, which need to support hundreds or thousands of users, Realvideo is suitable. The Realvideo format is proprietary and could be treated as an interim format, to be replaced by MPEG-4 or IPV when players and encoders are cost effective.

Audio Standards

Uncompressed 48 kHz AES/EBU locked to video

AC-3 audio in transport streams – for distribution, integration and archive only.

MP-3 for voice overs, and networking of voice overs.

44.1 sampled AES/EBU audio from CD-ROMS
RealAudio for super low bitrate browsing using a modem.

Metadata Standard

In the case of NBC, OMAR is the current metadata standard. Our MAM systems will incorporate a Metadata standard that embraces and extends the current OMAR system, so that NBC business units can archive more data and a structure that will support sophisticated searches for material in the archive.

Production Metadata

Intelligence Video Index, Content Description Metadata (Dynamic Metadata Dictionary Structure and Contents)

When the draft SMPTE standards for "Dynamic Metadata Dictionary Structure" and "Dynamic Metadata Dictionary Contents" have been approved by SMPTE and have been evaluated by the various companies that have a vested interest in this issue as meeting those companies requirements, they will become the likely STANDARD for all motion imagery metadata systems.

SMEF Data Model v 1.5 or later
Metadata system is developed by the BBC and is meant to be the STANDARD for motion imagery systems metadata.

Production Database and Local OMAR database

The structure and topology of MAM databases are key aspects to the performance and security. Production databases include data information from automation systems, newsroom computers, nonlinear editing systems and ingest cataloging systems. In a production database, ample information may simply refer to the specific phase a production essence clip may be in. That information is usually irrelevant to an archive database. We therefore see the need for a filter between the production database and the archive database. Librarians or archivists would "shape" the data during its migration from a Production database to Archive database, to make the archive search functions quicker and more accurate. The Production data can also be filtered by Middleware; software components that understand different structures in databases. Middleware is addressed later in this document.

Proposed Requirements for a MAM database might include the following:

- Database can be expanded easily, depending on storage architecture limitation.
- Data Records have limitations, however may be hyperlinked to other records with special processing.
- Database will either reside on one system with failover capability to one or more other systems, or will be accessed in parallel from two systems simultaneously
- System outage time will be required for upgrades hardware and software additions and new version rollouts
- An on line backup methodology will be set in place to provide complete data backup.
- A "read only" replication system may be employed for redundancy and performance enhancement
- Depending on Network architecture, low resolution accessible from producer workstation

- Depending on network architecture, Target browse datarate- Less than 1 Mbps data rate with audio
- Preferred application: Web browser
- Time to ingest tape to digital will be minimized with conversion system
- Work flow will allow for "look ahead digital conversion"
- Small snippet of taped material will allow one to partially view without complete digital conversion
- Database will be accessible to the user except for specific maintenance outage periods

Middleware

The middleware is actually the link to and from central archive and business unit. The middleware layer is the filter of data and formats going to archive and the transcoding from archive.

Data Messaging and Database Connectivity

Generally, handling the data in a local MAM System means observing standard IT rules for networking and redundancy. The MAM system should allow each business unit to use metadata that is unique or proprietary, and filter out meaningless data before migrating the media and data to archive.

Data messaging between systems can be handled with protocol tools like CORBA (Common Object Request Broker Architecture) XML (Extended Markup Language), COM, ContentShare and several others. A MAM system should strive for a unified approach to data messaging and connectivity.

Archive Business Logic

Business rules dictate the value and security of a clip and each operating group in the content game have their own business rules. For example, if Dateline shoots an investigative piece, they restrict display of their footage, until they agree it is OK to distribute. The same goes for Nightly News, Access Hollywood and many other production groups.

So, how do we reasonably migrate data and video to the archive, without violating business rules of each operating group?

Basically, three types of data in the internet-broadcast/archive system are anticipated: High Resolution Essence, Low Resolution Essence, and Metadata. The data will be stored on several media, including Videotape, Datatape, Optical Disk, Staging Hard Disks, On-line Server Hard Disks.

The goal of a MAM system should be to provide the correct access to these data, cost effectively.

It is helpful to query the business rules of each production group:

1. How much material gets ingested or recorded in the facility, daily?
2. Weekdays or Weekends or 7/24 ?
3. How long do you need to keep material on-line, in the facility?
4. How many dubs do you make of incoming material?
5. Do you erase any media that you know will not be re-used?
6. Do you have rights to the essence used on air?
7. What happens to your media when you run out of shelf space?
8. How long do your produced stories take to completion?
9. Do you need to "Hotel" your editing seats, have several users in the same seat during a day or week?
10. When could your production groups make information about your essence available to insiders and outsiders?
11. When could you release essence rights to insiders and outsiders?

Answers to these questions determine workflow and timing to and from the archive.

Data Backup

While it may be anticipated that metadata will always remain on-line. Eventually, the metadata will need to be backed up to some permanent storage and taken to another location, for catastrophic disaster protection. When will we make backups and where will the data be stored?

Replication of essence will also be required, based on popularity and needs. Media will need to be checked out and checked in, if, for example, the Robotic Library's shelves are full and media must move to shelves inside or outside the facility.

Automation System

An automation system can provide cost benefits, especially when tied to a Digital Asset Management System. But, the talent involved in high quality production may be a resource necessary for a good product. Thus, automation cannot be employed throughout the enterprise, but discreetly, where it will enhance the speed of production and not detract from production values recognized by the viewer. Newer digital systems in the home may provide "packaging" services that preclude the need for traditional "TV Show" packaging. When such a system provides a platform for consumption that is acceptable to several millions of "view-sers" the resources of a media asset management system may be a requirement to satisfy the capabilities of the consumer system. Such systems could allow show "Clipping", viewing only parts of traditional shows that are attractive to the viewer, or interactive shows, that branch in realtime.

Automation during the distribution phase has provided clear cost benefits, through the reduction of staff required to handle distribution. New markets may be necessary to justify the cost benefits of a Media Asset Management system, if production is the activity focus of the system.

System Build

The media asset management system is a complex system with functional dependencies. Phases should be proposed, which build, test and measure cost benefits. Phased projects allow users to adopt new functions gradually, with a more complete understanding of the new processes. Finding the correct functions to justify the costs may require several distinct analyses of build out scenarios.

Basic Questions:

Business unit cost-justification

Business Capital Budgets: Should each business unit justify its needs for a media asset management system, and build one if it can be justified to management?

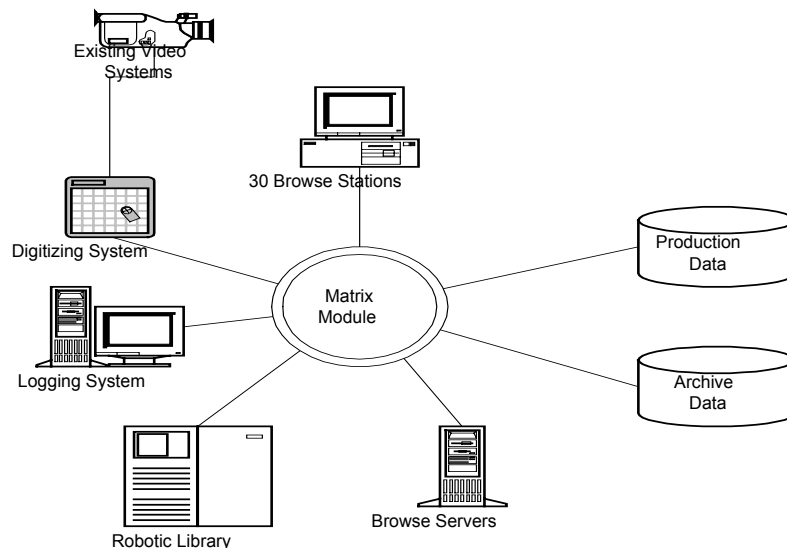
Management Needs for facilities

Do the efficiencies proven by a media asset management system provide compelling reasons for management to require that selected business units incorporate a media asset management system in their facilities.

the digital infrastructure developed for a browsing system, when high speed network and digital storage cost-performance improves to the levels required.

System Components and Vendors

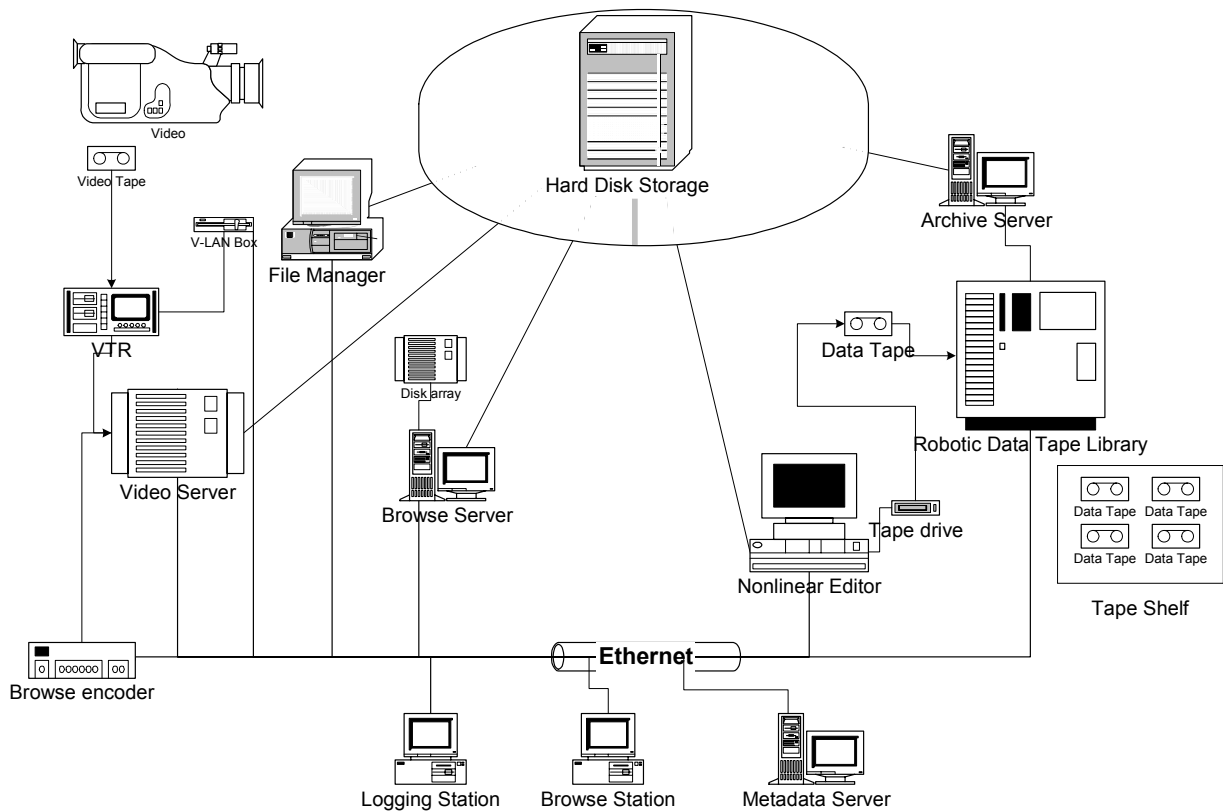
The following diagram illustrates a mid-term MAM solution that includes Ingesting, Producing and archiving high-resolution, browse/edit resolution and metadata. The basic MAM module does not include the Nonlinear editor, Storage Area Network and File Manager, or Video Server which are included in the mid-term solution.



A goal of the basic MAM module is to augment existing analog video and audio systems with advanced digital management techniques, which would support **sharing material, creating new products and streamlining the workflow.**

The MAM module would interface with text based search tools, a powerful database and robust hardware currently operating in News Archives. The MAM system would introduce infrastructure to handle digital essence, and improve the functions of searching.

The higher resolution analog video and audio systems in TV production can be replaced with



Conclusion

Each of us needs to understand the changes that technology will bring to our business and consider the need for a unified approach to Media Asset Management. The next logical step is to test the technology and vendors, develop firm cost benefits modeled from the test production system and proceed with a pilot build if business conditions warrant. Internet and interactive divisions should provide further impetus for the effort.

MIDDLEWARE: A KEY COMPONENT FOR BROADBAND

Dr. Ken Morse, Azita Manson

PowerTV, Inc.

Abstract

There are many discussions in the digital interactive cable television community regarding middleware. Debate rages over requirements, technology choices and even terminology and definition. This paper presents an overview of the design decisions and engineering challenges faced during the development of the pTV Software Platform by PowerTV, Inc. While the paper's focus is primarily client-based, it also addresses the approach adopted on the server side of the equation.

INTRODUCTION

This paper presents an overview of the engineering challenges of the pTV Software Platform from inception to delivery and subsequent evolution through seven years of development by PowerTV, Inc. The hardware target for this software platform is interactive Digital Home Communications Terminals (DHCTs) of the Explorer 2000 class and above. Alternative software approaches have been followed by vendors addressing the DCT 2000 class DHCTs and is not in the scope of this paper.

Before discussing the software architecture it is important to review the target DHCT capabilities and network

infrastructure in which the DHCT resides.

THE DHCT

The Explorer® 2000 DHCT provides support for digital services and traditional analog services delivered through a hybrid fiber/coax (HFC) network. The first version of the 2000 shipped in 1998 and was equipped with a 54MHz RISC microprocessor, 4MB of FLASH memory for system software, 2MB of memory for MPEG decompression and graphics, and 8MB of DRAM for system and application use.

The DHCT is equipped with 64 and 256 ITU J.83 [1] Annex B QAM demodulation support and a service tuner enabling both analog and MPEG 2 digital channels to be tuned and displayed. A DAVIC 1.1 [2] compatible out-of-band (OOB) system, operating at 1.544 Mbps is included to enable instantaneous, IP-based, "real-time" two-way communications between the DHCT and the headend.

Digital services are secured using the PowerKEY® conditional access system provided through an on-board security microprocessor.

Local devices may be attached to the DHCT through the Universal Serial Bus (USB) interface or an optional Ethernet

10BaseT interface. Together these support the connection of devices ranging from printers and digital cameras to personal computers.

DIGITAL BROADBAND DELIVERY SYSTEM

The DHCT terminates the digital broadband delivery system (DBDS) that combines video, audio, and data content from a variety of sources and distributes them to the subscribers' home. Given the scope of this paper, it is important only to understand the network connectivity as viewed by the DHCT.

Figure 1 illustrates the communications channels between the DBDS and the DHCT.

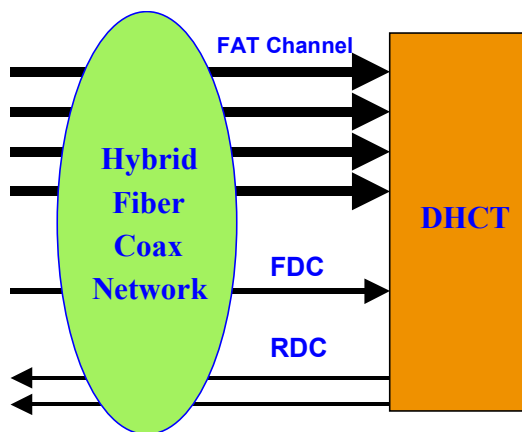


Figure 1 – Network Connectivity

- Forward Application Transport (FAT) channels. The DHCT can select any FAT channel by tuning to it.
- Forward Data Channel (FDC). The DHCT can always receive the FDC, even while tuned to analog services.

- Reverse Data Channels (RDC). The DHCT can only transmit in one RDC. However, more than one RDC may be defined per node for capacity reasons.

With the baseline understanding of the DHCT and the network connectivity a discussion of the software platform is presented.

SOFTWARE PLATFORM OVERVIEW

Once the DHCT platform and network were defined, it was possible to address the software architecture in the DHCT. The architecture follows a layered approach with clearly defined interfaces and responsibilities. Figure 2 provides a block diagram of the software components.

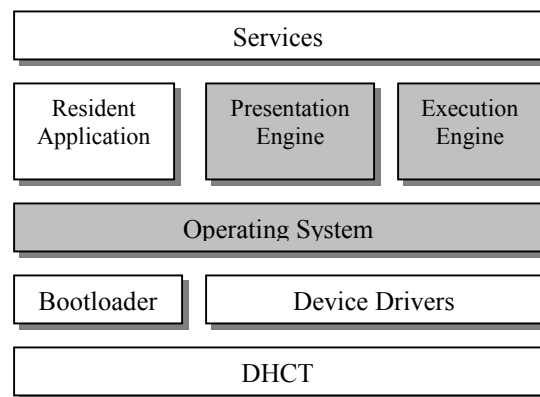


Figure 2 – Software Components

The items in gray make up the pTV Software Platform presented in this paper and are discussed in the following sections.

Bootloader

The bootloader is the gatekeeper of the DHCT and ensures the integrity of the system software stored in the

FLASH memory. It listens for communications from the headend to ascertain whether software updates are available for the particular hardware revision of the DHCT and also automatically recovers the correct version of software in the rare case of corruption in the DHCT image.

Device Drivers

Each instance of a DHCT (from the same or different manufacturers) may contain different components such as processor, memory, and custom semiconductors. To ensure that the software platform can operate across a range of DHCTs, a Hardware Abstraction Layer (HAL) is provided. This HAL is an Application Programming Interface (API) specification to which DHCT manufacturers author their device drivers to ensure compatibility with the software platform. It is important to note that the device driver implementation is the first layer of the software stack. Design and implementation decisions made at this level ripple through the entire software stack. Key implementation guidelines have to be provided at the device driver layer to address the efficient use of memory, interrupts, and thread prioritization. For example, since the DHCT deals with high bandwidth data flows, it is important to minimize the number of times a piece of data is copied as it traverses the software stack. In addition, information, such as diagnostics that may be required at higher levels in the software stack, is often generated within the device drivers and must be exported.

OPERATING SYSTEM COMPONENTS.

The Operating System (OS) provides the foundation on which all other software components are built. For this reason, it has three clear non-negotiable requirements.

Make the DHCT work

The DHCT must be robust and reliable as it is delivering core television services to the subscriber. Hence, the operating system must be designed to deal with the high data processing requirements in conjunction with the constrained memory environment while maintaining a high quality of service. In addition, the local interfaces on the box must be supported, such as USB and Ethernet.

Make the network work

The DHCT is the termination point in the network and as such the OS must support all the necessary protocols for interfacing with the rest of the network. This includes, but is not limited to, conditional access, system information table management, session management, carousel delivery and network management.

Provide a software platform

Since all other software components are built on top of the OS, it must provide an extensive API that provides access to all the functionality contained within the DHCT. For multiple software components to concurrently operate against the OS in a constrained DHCT, it must provide a comprehensive resource management environment enabling

policy decisions to be made by higher-level components.

Core Services

The core OS services are used by all software components in the system including all other OS components. A real-time, pre-emptive and multi-tasking kernel is required to ensure the deterministic operation of the DHCT when dealing with high bandwidth data flows. The kernel is optimized to manage low context switch times and low latency response to interrupts. In some DHCT designs there can be upwards of two hundred different interrupt sources – clearly identifying the complex design within a DHCT. In addition, several kernel modifications were made by the team based on the requirements of higher-level software components, such as Java to ensure efficient operation of the overall system.

Many DHCTs, including Explorer 2000, do not include a Memory Management Unit (MMU). This makes the task of ensuring the integrity and availability of memory a larger problem. The OS addresses this through a range of memory management processes applied to different types of memory in the DHCT – application, graphics, MPEG – in addition to algorithms to minimize memory fragmentation. As part of the resource management architecture, different software components may be asked to yield memory resources when memory availability becomes critical. This implementation maximizes the number of concurrent software tasks that may be operating on the DHCT.

It is important to note that the services in the OS must be multi-thread safe and re-entrant as necessary.

Multimedia Services

Another often-misused term is Multimedia and in the scope of the support from the OS this incorporates video, audio and graphics display. Video services are provided at the MPEG transport level in the OS, other video format support may be built in higher-level software components. Similarly, MPEG and Dolby AC-3 audio streams are managed as part of the MPEG transport stream. Additional support is provided for PCM and ADPCM audio sources in the OS.

Graphics functionality is split into three layers. basic graphics management provides access to the frame buffer(s) of the DHCT and supports multiple color depths basic graphics primitives in addition to pixel operations. These constructs are then used by the Window Manager to interface with the resource management architecture to allow multiple software components to share the television display concurrently. This includes combining graphics and the video plane of the DHCT.

The Window Manager design is the result of rationalizing the requirements presented by a diverse set of needs including higher-level software components such as Java and the Abstract Windowing Toolkit (AWT).

Finally, the OS provides a Widget Manager component enabling consistent use of Widgets throughout the software platform. This aids consistency of look and feel in the higher-level software components. Once again requirements on the Widget Manager included the Java AWT and widgets used in HTML pages such as buttons, edit fields, and forms.

Network Services

The DHCT is a network device and the OS must carry a set of core networking protocols. Some of these are generic IP-based protocols; others are specific to the cable television environment.

In the system previously described a DAVIC out-of-band implementation is employed. The OS carries the DAVIC signaling implementation and is responsible for managing the connectivity to the cable headend. This is a critical area affecting Quality of Service (QoS) in the DHCT. Extreme failures, such as power outages, require managed sign-on to the network to avoid congestion and overload. An indication of the current operational state of the network is required by all higher-level software components. The OS works behind the scenes to manage the integrity of the network connections and resources for those components.

Certain services such as Video On-Demand (VOD) require sessions to be put in place across the network and maintained – and on occasion, re-negotiated. Once again, the OS manages these connections for the higher-level components.

IP-based services are supported in the environment and an optimized TCP/IP network stack is in place to manage the IP packet flow. Modifications have been made to support multiple interfaces to the stack. This enables in-band IP support over the FAT channel and local routing of packets to other physical interfaces in the box.

A core component of any integrated design is network management. The OS

provides an extensive SNMP agent supporting a range of diagnostic information that may be queried over the network. It is also extensible and allows higher-level software components to dynamically install new Management Information Bases (MIBs) to report on items as diverse as QoS to subscriber patterns. The use of SNMP in the system greatly reduces the impact of introducing new services and components into an existing system.

Given the inherent broadcast nature of a cable television system, the digital broadband delivery system provides a Broadcast File System (BFS) for delivering data and objects to the set-top. This is a broadcast, carousel-based design and the OS maintains the available file list and retrieves objects as requested by higher-level components.

Television Services

Television services are the final high-level category in the feature set supported by the OS. These include the ability to acquire System Information (SI) tables from the network and tune to A/V programs. Once again, extensive QoS support is implemented in the case of SI outage on the network. The SI database is held in the DHCT and updated as new tables are delivered over the network.

Conditional Access (CA) is not limited to traditional television services and may be applied to any service available on the DHCT. CA may also be used for signing and authenticating software components delivered over the network to the DHCT.

Since traditional SI is limited to video/audio programming selections the

OS architecture was extended to support the dynamic selection of applications from remote sources, typically the BFS. On the client side, the Service Access Manager (SAM) provides the functions necessary to load, register, and launch an application on the DHCT. The API provides a method for loading an application into the DHCT memory, activating the application, and managing the application after it is loaded

RESIDENT APPLICATION

The Resident Application (RA) is built on top of the OS and provides the core digital television services to the subscriber. These include:

- Navigator
- Interactive Program Guide
- Impulse Pay-Per-View
- General Settings
- Emergency Alert System (EAS)
- Diagnostics
- Digital Music Service
- VCR Commander Support
- Virtual Channels

The RA utilizes the services provided by the OS to implement policy decisions for the set-top regarding service activation, suspension, and removal. The OS itself does not enforce policy; it provides the services through which another party may implement policy.

Resident Applications are available from Pioneer Digital Technologies and

Scientific-Atlanta for the platform described in this paper.

THIRD PARTY DEVELOPMENT

Given the consistent and open API of the OS, a range of third party developers have developed and deployed applications and services against the core platform. Each deployed application must go through a certification phase to ensure it is interoperable within the end-end system. Over one hundred companies are working against the OS today creating new and innovative services on the DHCT.

ALTERNATIVE SERVICE AUTHORING APPROACHES

While C-based services have been developed and deployed by a wide range of companies against the OS, it was clear early on that alternative authoring environments would have a place in the DHCT over time. Any authoring environment will have advantages and disadvantages. In deciding on alternative authoring approaches it is necessary to understand the benefits and deficits of C-based services. These are outlined below:

- + Performance
- + Footprint (for a given service)
- + Network utilization (self-contained)
- + Robustness
- Complexity of authoring

- Not cross-platform (processor-specific)
- Certification process

Alternative authoring approaches need to address some of the shortcomings of C-based services while not sacrificing some of the key advantages provided by the native approach.

As part of an investigation into alternative service authoring techniques in 1996, an existing public domain web browser was ported to the operating system outlined earlier. At this time the technologies involved were simple by today's standards; HTML 2.0 [8] for layout, GIF87 image format, and HTTP 1.0 for object requests. To validate the applicability of the browser environment an existing service, Video On-Demand, was modified, along with the browser, and tested on a real-world cable system. The modifications consisted of extending HTTP to provide control of the video server trick modes and adding several tags to HTML 2.0 [8] to enable the overlay of graphics and placement of MPEG video. It was clear that such a content authored approach with extensions for object request and control could form the basis of a future environment.

At about the same time, Sun Microsystems released the first version of its Java programming language in the form of a development kit, or JDK. By working closely with Sun, we were able to gain access to the source code of the Java environment and port it against the current software platform. This activity and the resulting work is detailed in the section, Execution Engine.

PRESENTATION ENGINE

Given the initial prototyping work on an HTML 2.0 [8] based browser, outlined above, a number of issues would have to be addressed if an HTML-leveraged solution was going to operate in the constraints of a home communications terminal. The approach decided upon was to design and develop a software component for presenting content driven services on-screen and communicating with known servers using modified data representations on HTTP. For this development, the term presentation engine (PE) was defined for the resulting product.

The requirements defined for the presentation engine are presented in the following paragraphs along with a description of the engineering challenges addressed in bringing the product to market.

Integrated into environment

The success of a platform is based on how well integrated the design is. With this in mind, the presentation engine must leverage the resource management architecture provided by the core platform and follow other good-citizen style guidelines. This also includes integrating with the existing Resident Application infrastructure.

Standards-based, where applicable

The goal of the development was to leverage existing standards as needed and extend to support platform-specific capabilities that were not addressed by the existing standard.

Meet footprint targets of DHCTs

Memory constraints in a DHCT are tight, typically 8MB to contain network buffers, platform state and IPG data with the remainder available for services and the use of the presentation engine state and the service content. The presentation engine itself would live in FLASH memory as part of the core platform.

Maximise leverage of existing content

It was clear that in the near-term there would be a lot of available content authored for the larger market, i.e., the Internet, and that there would be distinct advantages if this content could be leveraged into the platform in the early stages.

Enable platform capabilities

The DHCT provides a set of integrated features unlike other environments, such as Personal Computers. The presentation engine should provide access to these features.

Authoring tool availability

One of the complaints leveled against development in a standard procedural programming language is the complexity of the development. With this in mind it was important to ensure that the necessary tool chains were in place for a successful presentation engine product.

Robust and reliable operation

This is a television product and must be operational 24 hours a day, 7 days a week, 52 weeks a year - this cannot be over-communicated. The television experience is solid and dependable. The introduction of new services and technologies must not affect this in any

way. It was clear from the start of development that ensuring this requirement is met would be the toughest by far. Existing web-leveraged environments are either unable to meet this criteria, e.g. personal computer, or do not address the memory and processing environment of a DHCT.

PRESENTATION ENGINE DESIGN

From the previously stated set of requirements the design of a PE was undertaken. This section identifies the design choices made.

Approach

The decision was made early on to focus on what is meant by an engine. Even though the majority of content presentation solutions are based around some form of browser, the design team decided to architect the PE along the lines of the OS; i.e., a set of services that can be manipulated by higher-level components through well-defined interfaces. This approach ensures that resource management issues (such as multiple use instances) for the PE are addressed in the core design.

Content Support

Based on previous experience with the HTML 2.0 [8] based browser and the explosive growth of the World Wide Web in 1996 it was decided to implement along the lines of an HTML-based solution. This ensured that the availability of authoring tools would not be a problem. In addition the free-form nature of HTML page definition lends itself to manipulation to fit the television display.

Given the decision to utilize HTML the next question arose of what version of the specification to implement. This was decided based on the memory profiles projected in the DHCT arena over the initial product lifespan. Based on assumptions of 8MB DRAM platforms, HTML 3.2 [8] was chosen. HTML 3.2 [8] includes feature sets not normally supported on constrained memory devices such as frames and forms but given the expected utilization of television services (such as t-commerce and walled gardens) it was decided to support this feature set. The use of frames required focus on core navigation issues such as movement between different frames.

HTML pages can include a range of different objects beyond the HTML page description. These fall into two categories; MIME types, and scripting.

With a focus on leveraging existing tools and content within the constraints of the DHCT, the MIME types supported include:

- GIF
- JPEG
- AIFF
- WAV
- Text

These MIME types are supportable across the range of DHCTs deployed. Additional MIME types such as Macromedia Flash are supportable on more capable DHCTs.

HTML provides a very static environment and given the non-static

nature of the television environment it was decided that interactivity must be provided through the addition of scripting. ECMAScript was chosen due to its prevalence in the developer community and historical use on the World Wide Web. Traditional implementations of ECMAScript assumed an environment with extensive memory and memory management in place. Modifying ECMAScript to operate within the constraints of the DHCT environment provided to be an extensive task to ensure that no memory leaks were present and that ECMAScripts could run robustly.

Content Translation

HTML content may be authored for arbitrary page sizes; it does not have to conform to a given constraint since existing desktop browsers support scrolling in two dimensions. It was deemed that asking a subscriber to deal with two-dimensional scrolling in a television environment would be unacceptable. To solve this problem, all pages would need to be scaled to fit the available horizontal resolution of the window on-screen. This operation potentially involves extensive image scaling and table re-calculation, depending on the content. The DHCT is not equipped with the necessary processor bandwidth to support such tasks adequately and so the solution was to create a proxy/transcoder component that resides in the network. This proxy/transcoder intercepts DHCT requests, acquires the page and associated assets and scales them accordingly, passing the results to the DHCT (and to a local cache to maximize performance of subsequent requests).

Another key design point was to ensure that the proxy/transcoder did not have to scale based on the number of attached clients. The proxy/transcoder could be instead placed at the Internet Point-Of-Presence and scale to the content pipe, a much simpler task, resulting in tremendous cost savings to the operator.

In addition, the proxy/transcoder was designed to operate in an offline mode enabling “walled gardens” to be automatically created from existing content.

Performance

After robustness, most of the development activity centered around performance. This was classified in two primary areas:

- On-screen display
- Network performance

On-screen drawing performance was the key metric since this is what the subscriber sees and determines whether the product is adopted or not. Standard techniques such as double buffering provided difficult at high graphics resolution given the limited memory available for graphics buffers. This resulted in alternative drawing algorithms being implemented. In two-way systems, the Proxy/Transcoder could be utilized to ensure content delivered to the DHCT was scaled appropriately, if necessary, resulting in a performance improvement in perceived drawing speed.

A caching subsystem was implemented to handle pages and assets (in both compressed and decompressed

forms). The caching system increased the availability of content and therefore reduced perceived subscriber access time. Through HTML extensions it is also possible to assign a caching priority to pages and assets. In conjunction, a scheme for pre-fetching was introduced. These two features, when used in combination, are useful in “walled garden” environments to ensure fast access to frequently used assets and offers a degree of deterministic content control in the client.

Finally, the PE can also receive content delivered through other mechanisms than HTTP. These include the BFS, MPEG private sections and the Vertical Blanking Interval (VBI). Such broadcast asset delivery mechanisms can be used to minimize un-necessary upstream traffic on the network.

Secure Connections

Given the expected use of t-commerce on the platform a secure communication mechanism was required. Once again, the standard of the Web, Secure Sockets Layer (SSL) was adopted. SSL is supported in two versions (2.0 and 3.0) with a range of different strength ciphers. The Presentation Engine design enables selective loading of ciphers and protocols to address the memory constraints of current generation DHCTs.

While the proxy/transcoder can provide translation facilities for non-secure pages it cannot translate content contained on secure web pages (using SSL). Two approaches were considered to address this. One approach is to introduce a proxy that terminates the SSL connection at the headend. In-the-

clear page content can then be transcoded and then a separate proxy may be utilized to form a secure connection with the DHCT. This obviously breaks the end-end security from the merchant to the subscriber and introduces legal issues that make such a solution undesirable. The solution employed is to provide client-side re-layout of the pages in the case of secure connections. This has the downside of potentially affecting performance but does ensure the integrity of the end-end security.

ATVEF

The Advanced Television Enhancement Forum (ATVEF) specification provides an environment for program synchronous content delivery. The content may either be included with the video/audio content or stored separately on a remote server. ATVEF support is designed into the pTV Software Platform and requires no specific support from the PE beyond supporting the *tv:* URL. As intended, the PE merely acts as a service to present the necessary ATVEF content on-screen.

Status

The PE is part of the pTV Software Platform and has been shipped in over seven million DHCTs to date. The memory footprint is configurable from 700KB up to 1.4MB.

EXECUTION ENGINE

The Presentation Engine is designed to display web content documents to the users. This model yet useful and simple to implement cannot accommodate for all the needs of an interactive

application, so there is need for a procedural programming language, which is easy and fast to implement. The industry has chosen the Personal Java for the Execution Engine as it is defined by the DVB-MHP and OCAP.

PowerTV and Sun Microsystem have collaborated to create a PJava implementation that is optimized to run in an interactive two-way cable network.

PJava is designed for a constrained network connected devices. The small memory footprint of PowerTV Operating System combined with Personal Java enables this design to run on the low-cost, two-way capable DHCTs. Write Once, Run Anywhere™ design of PJava allows for the fast and low cost development of applications. With PJava design, a prototype may be built and debugged on computer and deployed on the DHCTs.

pTV Software Platform is comprised of the Java Virtual Machine and set of standard class libraries defined by the PJava, version 1.1.

By using the common underlying resource management, pTV Software Platform has been able to achieve an environment, which allows for both presentation engine and execution engine to coexist harmoniously.

Some of the enhancements made to the pTV Software Platform, while implementing the JVM, was to increase the support for additional threads in the kernel. Memory management has also been modified to accommodate for the JVM memory requirements. The existing Window Manager has been modified to fully support the AWT user interface model.

In future, pTV Software Platform shall implement support for the JavaTV APIs. This effort is minimal due to the full support of all core TV functionality in the existing pTV Software Platform.

FUTURES

New Functionality

With the current status presented the question is “what of the future?” As ever, the key drivers in the interactive cable arena are the trends in DHCT design and deployment mechanisms. New functionality is being introduced into DHCTs including PVR, Home Gateway, etc. These features will drive a necessary extension of the middleware interfaces on the DHCT.

OpenCable

Perhaps a larger influence is the Federal Communications Committee (FCC) mandate for cable operators to give subscribers the opportunity to acquire their DHCT through retail outlets. CableLabs® has created a sub-committee to address this requirement. In addition to defining the hardware platform requirements and conditional access interfaces, OpenCable® has a working group defining the application platform for the retail DHCTs.

OPENCABLE APPLICATION PLATFORM (OCAP)

The OCAP specification defines the Application Programming Interface for an OpenCable™ retail device. The Middleware API, defined by OCAP, is independent of the underlying hardware. Furthermore, the OCAP middleware layer provides an abstraction layer to the

set-top box's operating system that manages the underlying hardware's resources.

The OCAP specification is based on the DVB-MHP specification. OCAP specification describes in details the deviation from DVB-MHP, where applicable, in addition to the new components definition. OCAP specification has been designed to focus on the needs of the North America's cable market.

The OCAP specification is intended to create an open environment for application developers to write applications, which are platform independent.

DVB-MHP

DVB-MHP's architecture is based on the DVB-J platform, which includes a Java Virtual Machine as defined by Sun Microsystems. DVB-MHP is comprised of three layers, applications, system software or middleware, and the resources on the set-top box.

OCAP Basic Architecture

The OCAP architecture resembles closely to the architecture illustrated in Figure 2. The OCAP architecture is comprised of two major components: the Presentation Engine and the Execution Engine.

What is OCAP Presentation Engine?

The Presentation Engine supports a set of declarative applications. The Presentation Engine is required to support the content formats for markup (HTML, XHTML) [8], cascading styling sheet (CSS) [8], plus Application Programming Interfaces (DOM) [8].

What is OCAP Execution Engine?

The Execution Engine is based on the Java Virtual Machine and it is comprised of set of APIs defined by PersonalJava Application Environment (PJAE) [3], Java TV API Specification [4], Java Media Framework (JMF) Specification [5], Home Audio/Video Interoperability (HAVi) Architecture Specification [6], and Digital Audio-Visual Council (DAVIC) Specification Part 9 [7].

OCAP Minimum Platform Requirements

The minimal set of required device resolutions that OCAP terminals shall support is 640x480 (square pixels) for Background, video, and graphics. These resolutions shall be supported for display aspect ratios of 4:3 and 16:9.

The minimum processor capability has been defined as 200 Mega Instructions Per Second (MIPs).

The memory requirement is 64 M Bytes of RAM and 16M Bytes of ROM. The NVRAM shall provide at least 16K Bits maximum storage.

SUMMARY

In closing, it is clear that middleware plays an active role in the increasing adoption of interactive television systems. While it is important to have consistent interfaces to which services can be authored it is also clear that an optimized implementation of the software stack providing these interfaces is necessary. This ensures performance, maintainability, extensibility and time to market – all key requirements to future success of the industry. While different middleware standards have been implemented over the past few years

there is hope for a common interface through activities in the retail space such as the OpenCable initiative. However, for current non-retail systems, the challenges of building an optimized middleware stack remain and integrated platforms present the best path forward.

ACKNOWLEDGEMENTS

The authors would like to thank the entire team at PowerTV. Without them and their developments this paper would not exist.

REFERENCES

- [1] ITU Recommendation J.83 (04/97) - Digital multi-programme systems for television, sound and data services for cable distribution; Annex B
- [2] Digital Audio Visual Council™ (DAVIC), version 1.1
- [3] Sun Microsystems, PersonalJava Application Environment Specification Version 1.2a:
<http://java.sun.com/products/personaljava/>
- [4] Java TV API Specification;
<http://java.sun.com/products/javatv/>
- [5] Java Media Framework, Version 1.0, May 11, 1998; Sun Microsystems Java Media Framework Specification
<http://java.sun.com/products/java-media/jmf/1.0/apidocs/packages.html>
- [6] Home Audio/Video Interoperability (HAVi) Architecture Specification;
<http://www.havi.org>
- [7] Digital Audio Visual Council™ (DAVIC), version 1.4.1 Part 9.

[8] www.w3.org

MIGRATION TO FULL DIGITAL CHANNEL LOADING ON A CABLE SYSTEM

Marc Ryba

Motorola Broadband Communications Sector

ABSTRACT

Present day cable systems run a mix of both analog and digital signals. As digital services evolve, HDTV becomes reality, and competition from fully digitized direct satellite services increases, migration toward a full digital spectrum is inevitable. Moving to a full digital multiplex raises some concerns with regard to distribution and signal quality. Digital signals produce noise-like distortions that extend throughout the spectrum. Full digital channel loading will decrease the measured C/N (E_s/N_o) in the QAM channels. This paper will present the results of lab tests to determine optimum levels for digital signal distribution and the effect of a full digital multiplex on a modern cable plant. Digital signal quality issues are also investigated.

BACKGROUND

With the rollout of HDTV and competition from fully digital home satellite systems, increasing the number of digital channels in a cable system is the next step in digital deployment. Increasing digital services will attract new subscribers, provide higher quality services and minimize existing subscriber turnover.

Typical cable systems are currently running a mixture of both analog and digital channels with the digital signals appended to the higher frequency edge. Due to distribution system non-linearities, it has been shown that digital distortion produces a noise-like spectrum that extends into the lower adjacent analog channels [1]. This phenomenon is known as Composite Intermodulation Noise (CIN) and has been described in earlier publications [2].

Systems evolving to a full digital multiplex will also need to concern themselves with the effects of CIN throughout the distribution system. Full digital loading will produce noise-like distortions that extend throughout the spectrum. Characterization of digital second order intermodulation products shows that they always fall outside the frequency range of the digital signals that generate the products. Third order intermodulation products fall into several distinct categories. These distortion products have components that fall outside the frequency range of the source signals and also components that fall within the signal bandwidth. Second and third order distortion products (those falling outside the generating signal's bandwidth) are statistical in nature, additive, and uncorrelated in time with respect to the active signals [3]. This will decrease the

measured C/N or Es/No in the digital channels, with the greatest degradation occurring near the center and edges of the frequency range. The effect will be proportional to the number of digital channels present and their overall power level [1].

A third order distortion product that falls within the bandwidth of the signal that generates the product is classified as a cross-modulation product and has a multiplicative effect. These products are correlated with the generating signal and cannot be measured using spectral analysis [3]. Therefore, the measured Es/No values presented in this paper exclude the cross modulation effect.

Operating the digital signals at reduced power levels through the distribution system minimizes the effects of intermodulation products. A balance between noise and distortion products must be achieved in order to optimize Bit Error Rate (BER) performance. Tests were conducted with the objective of determining optimum levels for digital signal distribution, consistent with avoiding fiber optic clipping and minimizing CIN. Es/No ratios and digital signal quality were measured as a function of frequency throughout the spectrum. Digital signal quality in the test channels was gauged according to Bit Error Rate (BER) and Modulation Error Ratio (MER).

TEST SYSTEM DESCRIPTION

The system configuration for the tests consisted of a 133 digital channel multiplex in the 55-860 MHz range. Tests were conducted on 9 channels spaced throughout the spectrum. A block diagram of the test system is shown in Figure 1. The digital multiplex was generated using a platform of 49 independent digital signals received from various satellites and transcoded to 64 QAM via a Headend consisting of Motorola IRT1000 and IRT2000 Integrated Receiver/Transcoders.

To generate the full multiplex, the spectrum was broken up into 4 distinct frequency groups as follows: 55.25 to 97.25 MHz (Group 1), 97.25 to 349.225 MHz (Group 2), 355.25 to 601.25 MHz (Group 3), and 607.25 to 853.25 MHz (Group 4). Group 1 channels were received from 7 IRT's and independently upconverted using Motorola C6U upconverters. The IF outputs of 42 IRT's were split with each leg feeding two separate C6U upconverters that covered groups 3 and 4. Group 4's RF spectrum was split with one leg down-converted using a mixer to generate the signals in Group 2. Group 2 and 3 were de-correlated from the original signals using two fiber optic links with delays of 1 km and 2 km, respectively. All 4 groups were then combined to form a contiguous spectrum of uncorrelated digital channels.

A notch filter was used to reject noise contributions prior to the distribution system for each channel under test. The filter was inserted after the combined output of a frequency group and prior to the

group's introduction into the combined multiplex. This provided better than 50 dB of noise rejection in the channel under test. The channel that originally occupied the space in the multiplex was turned off and a 64/256 QAM signal transmitting a Pseudo-Random Bit Stream (PRBS) was inserted in order to make BER and MER measurements. Once the test on the channel was completed, the filters were removed and the original channel reinserted. Filters were available for all channels except EIA Channel 134. The Carrier to Noise ratio measured on this channel was better than 49 dBc and the measurement error due to the absence of the notch filter was assumed to be negligible.

Figure 2 Shows the HFC distribution system. The digital multiplex was input to a medium sized HFC distribution system. The optics consisted of an ALM-11 1310 nm laser transmitter, 20 kilometers of fiber, and an SG-2000 Fiber optic node. The RF system consisted of two Motorola BLE-87S/G GaAs line extenders and 21 taps. The RF distribution system was designed with 11.5 dB of tilt for 860 MHz spacing. The SG-2000 and Line Extenders were operated in manual mode with the AGC disabled. A characteristic of this particular distribution system displayed a slight flattening of the frequencies above 700 MHz originating from the SG-2000 node. This propagated through the distribution system to the 2nd line extender.

The level of the digital multiplex

was adjusted to provide an input level that is near optimal for the fiber optic transmitter. A power meter was used at the laser's RF test point to monitor AGC effects on laser drive level and provide an indication of modulation. The transmitter's AGC works on average power and maintains a constant drive level to the laser. Motorola's Headend Control Software (HCS) was used to monitor and control the transmitter's laser drive level, optical output power, and the laser Depth of Modulation (DOM). With AGC enabled, the average power at the laser test point varied according to the laser DOM between -23.5 to -25 dBm.

Digital C/N (E_s/N_o) ratios were measured using a HP89441 Vector Signal Analyzer. A tunable bandpass filter, centered on the test channel, was used for all measurements to eliminate saturation and distortions at the input to the test equipment. In-band ripple of the tunable channel filters did not exceed 0.5 dB in a 6 MHz bandwidth. The digital test signal was monitored after the second line extender throughout the test. This test point provided a signal that was greater than +25 dBmV per channel for accurate C/N measurements. Two Motorola DCT2000's, modified to run Broadcom QAMLink™ software, were used to monitor the Pre and Post FEC BER, and MER of the digital channel under observation. A constant signal level of 0 dBmV (± 2 dB) was maintained at the set top input throughout the test.

TEST RESULTS

Tests were performed with the laser's AGC enabled in both Preset and Set Modes. The AGC maintains the laser DOM in both factory set and user offset levels. Preset Mode allows the laser AGC to maintain the Depth of Modulation (DOM) at factory-default level. This level is optimized for a full load of 110 NTSC channels. Set Mode allows the DOM to be incremented or decremented from factory set levels by 0.25 dB increments. For a typical analog system, reducing the laser DOM by 1 dB degrades the C/N ratio by approximately the same level and improves intermodulation distortion products by approximately 2 dB and 1 dB for CTB and CSO, respectively. Test results for the different modes of operation are given below. All MER and short-term BER tests were performed using a PRBS source for both 64 and 256 QAM on all test channels.

Testing was also performed with a reduced channel loading of 91 digital channels. The amplifier drive level was adjusted to meet the same Carrier/Noise ratio and C/CTB and C/CSO ratios as with the full digital multiplex.

Preset Mode

Initial link analysis was performed with 133 CW carriers to adjust the distribution systems operating level with respect to a full analog spectrum. The number of NTSC channels specifies the input level (total power or power per channel) to

the laser transmitter. The nominal level specified for 133 NTSC channels is approximately 12.5 dBmV/Channel. The laser was operated in the preset-CW mode for performance testing which increases the drive level to the laser by 3 dB. Output of the distribution system was set to typical C/N, C/CTB and C/CSO operating levels when running a fully loaded modulated NTSC spectrum. Table 1 shows the performance of the HFC distribution system with full analog loading measured at the node and line extender 2 (BLE #2) outputs. The optimum output level of the distribution system was found to be 32.5 and 43.5 dBmV for 55.25 and 853.25 MHz, respectively, to achieve typical performance parameters.

For preset testing in video mode, the Laser Depth of Modulation (DOM) is held constant by the AGC at 0 dB and the Laser Drive Level at 5.75 dBm. The average power of each digital signal was set according to analog specifications described above. Initial tests were run to find the effects of amplifier cascade drive level changes on the distribution system. To find the optimal levels for running the RF cascade, the input level to the transmitter was held constant and the SG-2000 RF output level was varied around the optimum NTSC operating point of 32.5/43.5 dBmV. This optimum operating point represents the average power level that would be seen with a fully loaded NTSC system. Figure 3 shows the results of 64 QAM Es/No Ratio versus Frequency performance for different RF cascade operating levels.

Higher levels of distortion are seen on the higher frequency channels as the power level of the RF section is increased. This can be attributed to the positive tilt in the distribution system, which causes increased amplifier distortion at the higher frequency channels. The low Es/No ratio levels at the lower frequency edge can be attributed to second order intermodulation distortion.

Figure 3 shows that, when the amplifier section is run into compression, digital third order distortions are reduced approximately 1.5 dB by a 1 dB decrease in amplifier drive level. This decrease in the noise floor extends to an operating level approximately 2.5 dB above the distribution system's optimum setting of 32.5/43.5 dBmV. This effect is predominant in the higher frequency range. Below this level, a 1 dB decrease in output level shows less than a 1 dB decrease in Es/No ratio performance throughout the spectrum. Operating the RF section for improved distortion performance as opposed to optimum C/N ratio improves overall system performance. The recommended Es/No ratios are 28 dB and 34 dB for 64 QAM and 256 QAM, respectively. From the figure, the worst-case recorded Es/No ratio of 38.2 dB exceeds the minimum recommended Es/No ratio operating level for both modulation schemes.

Test data in Table 2 and Table 3 show MER and short term Pre-Reed-Solomon (Pre-RS) BER performance for both 64 and 256 QAM. Pre-RS

BER is measured after Trellis Decoding in the QAM decoder. Note that the Pre-RS bit error rates for all short term tests were well within acceptable levels of operation. Test data on Pre-RS BER performance for 64 QAM signals indicate no errors occurring in the test channels. Improvement should be seen in the 256 QAM Pre-RS BER rate as the time duration is increased as shown in long term BER testing. No burst errors were recorded and the Post-FEC BER for all tests was recorded at zero errors. MER values for both 64 and 256 QAM test channels are also reasonable. An improvement in MER for both 64 and 256 QAM is seen as the amplifier drive levels are reduced.

Measurements were made to determine the Es/No ratio degradation at various depths within the RF section under the optimum amplifier drive level. Figure 4 shows the Es/No ratio versus Frequency for the RF section of the distribution system. Table 4 shows test data for both 64 and 256 QAM MER. Data shows that digital signal degradation is minimal through the different amplifier stages as long as the bandwidth and minimal operating conditions for the distribution system are met.

Set Mode

Since the laser is optimized for 110 NTSC channels in preset video mode, the transmitter was changed to set-mode in order to control and change the Depth of Modulation.

Figure 5 shows the effect of changes in DOM on noise performance. The transmitter shows approximately a 1 to 1 dB change in E_s/N_0 ratio degradation versus a decrease in DOM. This demonstrates that the laser transmitter can effectively be run in the preset video mode and is virtually transparent to the digital signals with full digital loading. Table 5 shows 64 QAM MER measurements as a function of DOM. No advantage is gained by decreasing the DOM with full digital loading. It can be seen from Figure 3, Figure 4 and Figure 5 that the RF section of the distribution system is the major contributor to digital distortions in the link.

Laser Clipping

The Laser Drive Level (LDL) was adjusted to measure the system's laser clipping margin with full digital loading. RF amplifier compression was considered to be negligible for the increase in node level variations covered in this test. To measure this effect, the Pre-RS BER was monitored while the LDL was manually increased. Laser clipping produces burst errors in the digital signal that cause large bit error counts during testing. Clipping can also cause the digital demodulator to momentarily lose synchronization. MER is a time-averaged measurement and does not effectively capture periodic clipping events; therefore, BER testing must be used to capture such events.

Bit error rate (BER) measurements were made on the

test channels located throughout the digital multiplex. Short-term test results shown in the tables above demonstrated no clipping events. The LDL for 64 QAM was increased by 6 dB with no Pre-RS BER events recorded. A Pre-RS BER increase from $10E-9$ to $10E-7$ was recorded for 256 QAM signals at this same level. Since the Pre-RS BER is measured after Trellis decoding, some errors will be corrected and remain undetected.

Reduced Digital Loading

Reduced digital channel loading shows similar results. The decrease in the number of channels allows for the RF section to be run at a higher output level to achieve the same performance. Figure 6 and Table 6 show test data for a 91 digital signal multiplex.

Long Term BER Testing

Long-term BER tests were performed at random using the two laser modes. The tests were conducted using a PRBS 256 QAM signal located on the higher frequency channels. Table 7 presents long term bit error rate test data for different configurations of the distribution system. Note that the Pre-RS BER rate for all long term tests were on the order of $10E-9$ and well within acceptable levels of operation. No burst errors were recorded and the post-FEC BER for all tests was zero errors. The MER remains constant under the different loading conditions.

CONCLUSIONS

Cable plants will soon need to transition to a fully digital multiplex due to the advent of terrestrial HDTV and the demise of analog services along with increased competition from fully digitized satellite services. Transition to full digital operation may come sooner than expected.

Amplifier reduced rebuilds in cable distribution systems typically run at levels with a C/N ratio near 49 dB and C/CSO and C/CTB ratios greater than 53 dBc. Test data show that implementing a full digital spectrum can be realized on modern distribution systems as long as system bandwidth and performance specifications are met. Test results show that digital distortion can be regarded as an additional noise source that adds to system AWGN to produce an increase in the noise floor throughout the active spectrum. Test data show that this effect remains relatively flat in the spectrum as long as the RF section is operated within reasonable limits. Minimizing the generation of digital distortion in the distribution system will improve the Es/No ratio and optimize digital signal robustness. Test data show that operating the distribution amplifiers slightly below specified operating levels can reduce system noise due to digital distortion. The worst-case recorded Es/No ratio of 38.2 dB exceeds the minimum recommended Es/No ratio operating level for both 64 and 256 QAM signals.

Current amplifier stages rely on a peak power AGC that derives its

control information from an analog signal. The strategy of using peak power in the RF amplifier AGC section needs to move towards an average power type detector. This will provide stable levels due to temperature changes in the distribution system.

Transition to a full digital multiplex will soon become a necessity for cable operators.

ACKNOWLEDGMENTS

The author would like to thank the entire staff of Motorola BCS, Special Projects Laboratory, in particular Christopher Lynch and Robert Wanderer. Also special thanks to Charles Moore and Bob Lisher from Motorola, BCS for their assistance in completing this project.

REFERENCES

- [1] M. Ryba, J. Waltrich, "The Effect of Digital Signal Levels on Analog Channels in a Mixed Signal Multiplex", NCTA 2000 Technical Papers
- [2] J. Waltrich, "Distortion Produced by Digital Transmission in a Mixed Analog and Digital System", Communications Technology, April, 1993
- [3] D. Jaeger, "Uebertragung von hochratigen Datensignalen in Breitbandkommunikationsnetzen", Doctoral Thesis, University of Braunschweig, February, 1998

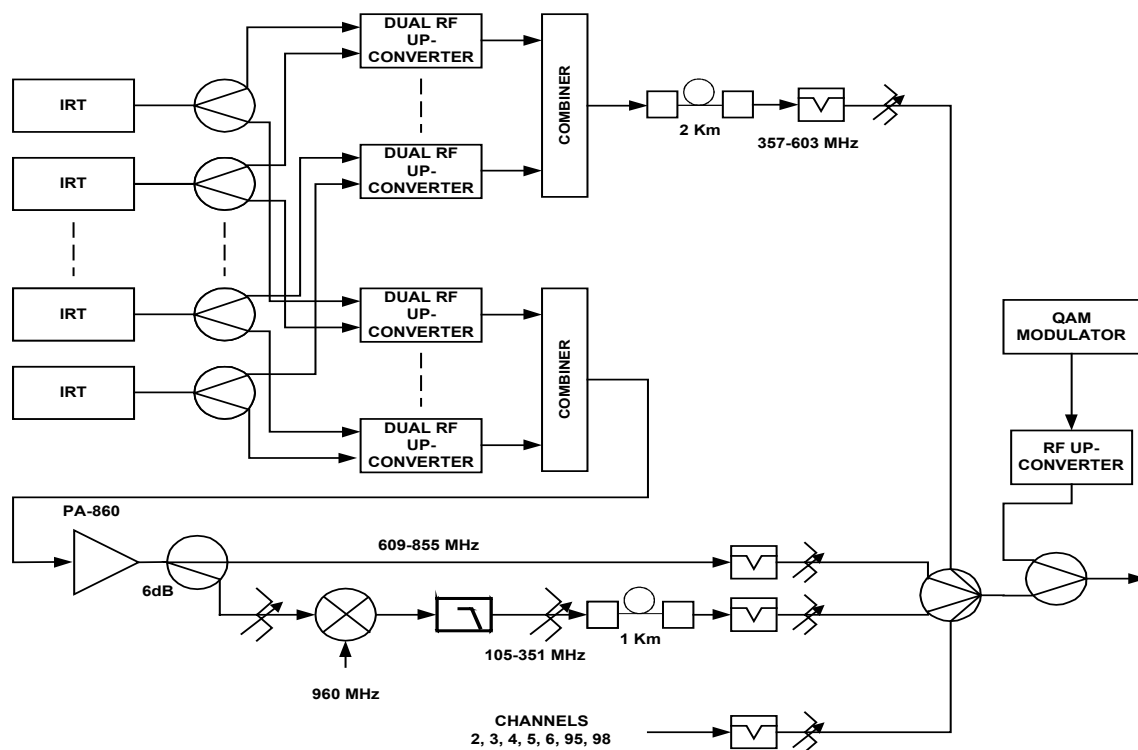


Figure 1 – Headend Configuration

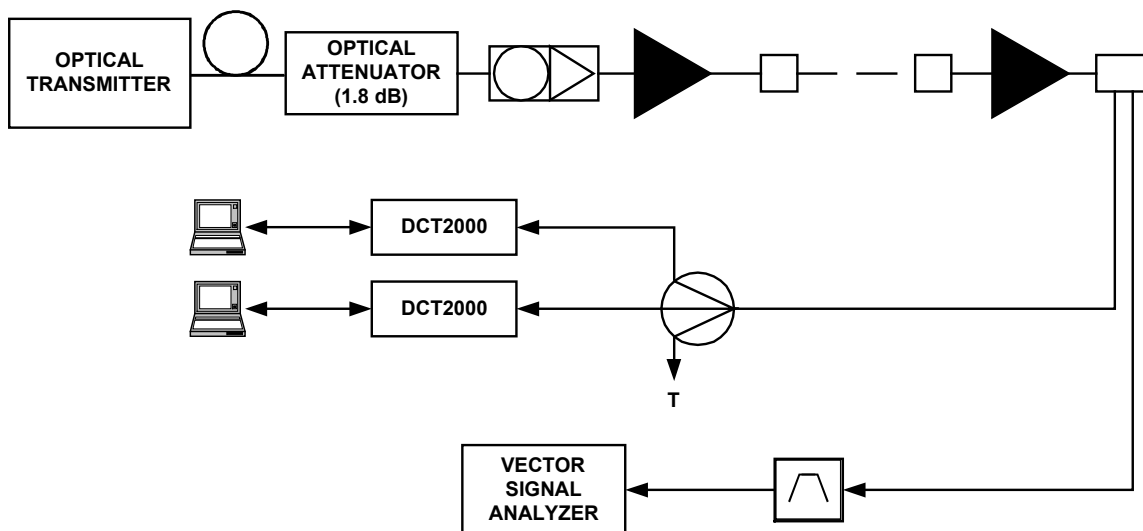


Figure 2 – HFC Configuration

Table 1 – 133 Channel CW Performance Summary

CW Performance						
Test Frequency (MHz)	C/N (dB)		C/CTB (dBc)		C/CSO (dBc)	
	Node	BLE #2	Node	BLE #2	Node	BLE #2
55.25	50.8	50.1	59.7	63.0	60.4	56.6
325.25	51.1	50.3	55.7	53.2	71.4	78.2
649.25	51.3	50.6	79.8	54.9	72.9	71.0
745.25	52.5	51.5	62.4	54.0	67.3	64.1
853.25	51.5	49.6	60.8	54.2	61.9	61.5

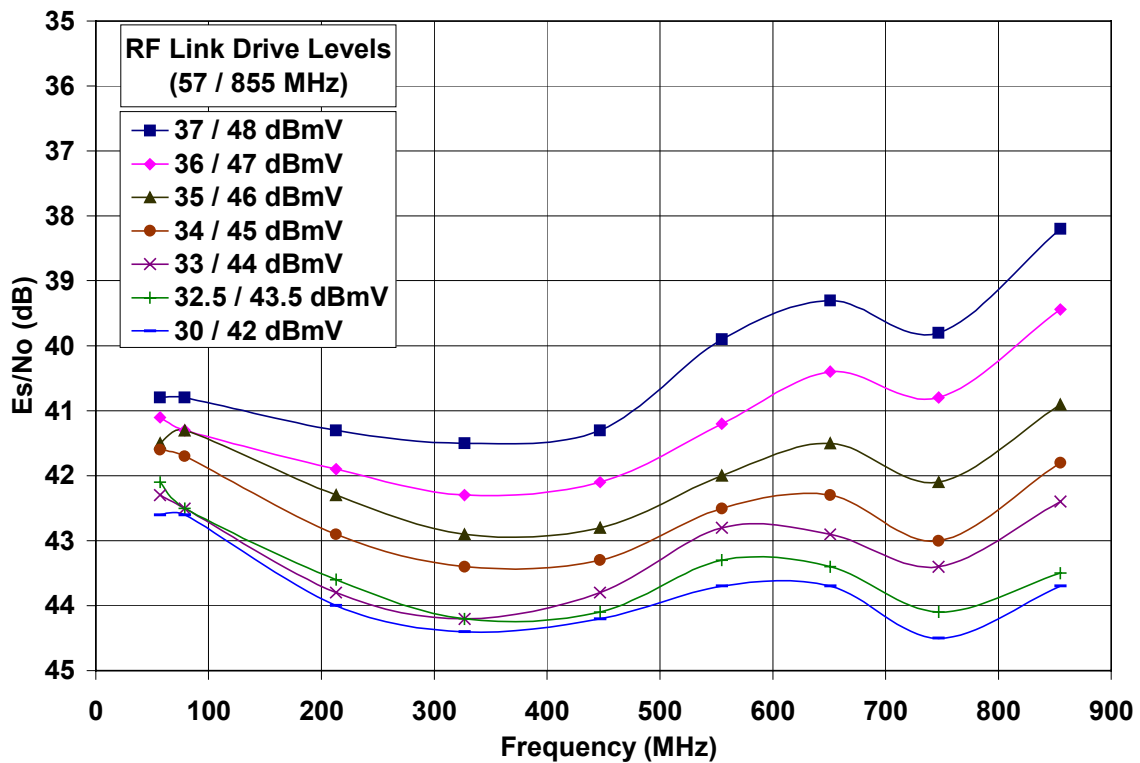


Figure 3 – SG-2000 RF Amplifier Output Level Variation for 64 QAM

Table 2 - 64 QAM Data for Various RF Amplifier Drive Levels

64 QAM - Laser Preset Video Mode														
Center Frequency (MHz)	SG-2000 RF Output Level 55 / 855 MHz (dBmV)													
	37 / 48		36 / 47		35 / 46		34 / 45		33 / 44		32.5 / 43.5		30 / 42	
	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER
57	34.3	0	34.2	0	34.4	0	34.4	0	34.5	0	34.6	0	34.6	0
79	34.4	0	34.6	0	34.4	0	34.6	0	34.7	0	34.7	0	34.7	0
213	34.1	0	34.3	0	34.7	0	34.6	0	34.9	0	34.8	0	34.9	0
327	34.3	0	34.5	0	34.6	0	34.6	0	34.7	0	34.8	0	34.8	0
447	34.6	0	34.7	0	34.8	0	34.8	0	35.0	0	34.9	0	34.9	0
555	34.1	0	34.6	0	34.6	0	34.9	0	35.0	0	34.9	0	35.0	0
651	33.9	0	34.4	0	34.6	0	34.8	0	34.8	0	35.0	0	34.9	0
747	34.2	0	34.6	0	34.8	0	34.9	0	35.0	0	35.2	0	35.1	0
855	33.6	0	33.9	0	34.5	0	34.5	0	34.8	0	34.9	0	34.8	0

Table 3 - 256 QAM Data for Various RF Amplifier Drive Levels

256 QAM - Laser Preset Video Mode														
Center Frequency (MHz)	SG-2000 RF Output Level 55 / 855 MHz (dBmV)													
	37 / 48		36 / 47		35 / 46		34 / 45		33 / 44		32.5 / 43.5		30 / 42	
	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER
57	33.1	2.1E-09	33.2	0	33.3	0	33.3	0	33.4	3.2E-09	33.4	4.2E-09	33.5	0
79	33.4	0	33.4	0	33.5	0	33.6	6.4E-09	33.6	5.7E-09	33.7	7.9E-09	33.8	1.2E-08
213	33.3	3.7E-08	33.4	0	33.5	5.0E-09	33.6	4.1E-09	33.7	0	33.6	4.8E-09	33.4	0
327	33.1	0	33.2	0	33.3	8.4E-09	33.3	1.5E-08	33.5	0	33.7	2.7E-08	33.6	0
447	33.4	0	33.5	0	33.6	1.1E-09	33.6	1.2E-09	33.7	0	33.5	5.6E-09	33.6	4.4E-09
555	33.2	0	33.5	0	33.6	2.9E-09	33.7	0	33.8	0	33.8	0	33.9	0
651	33.0	4.7E-09	33.3	4.2E-09	33.6	0	33.7	5.7E-09	33.8	0	33.9	0	33.9	0
747	33.3	1.3E-09	33.7	1.2E-08	33.9	0	33.9	6.9E-09	34.2	0	34.2	0	34.1	6.9E-09
855	32.6	5.1E-08	33.0	1.7E-08	33.4	3.7E-08	33.5	1.1E-08	33.6	0	33.7	4.8E-09	33.7	0

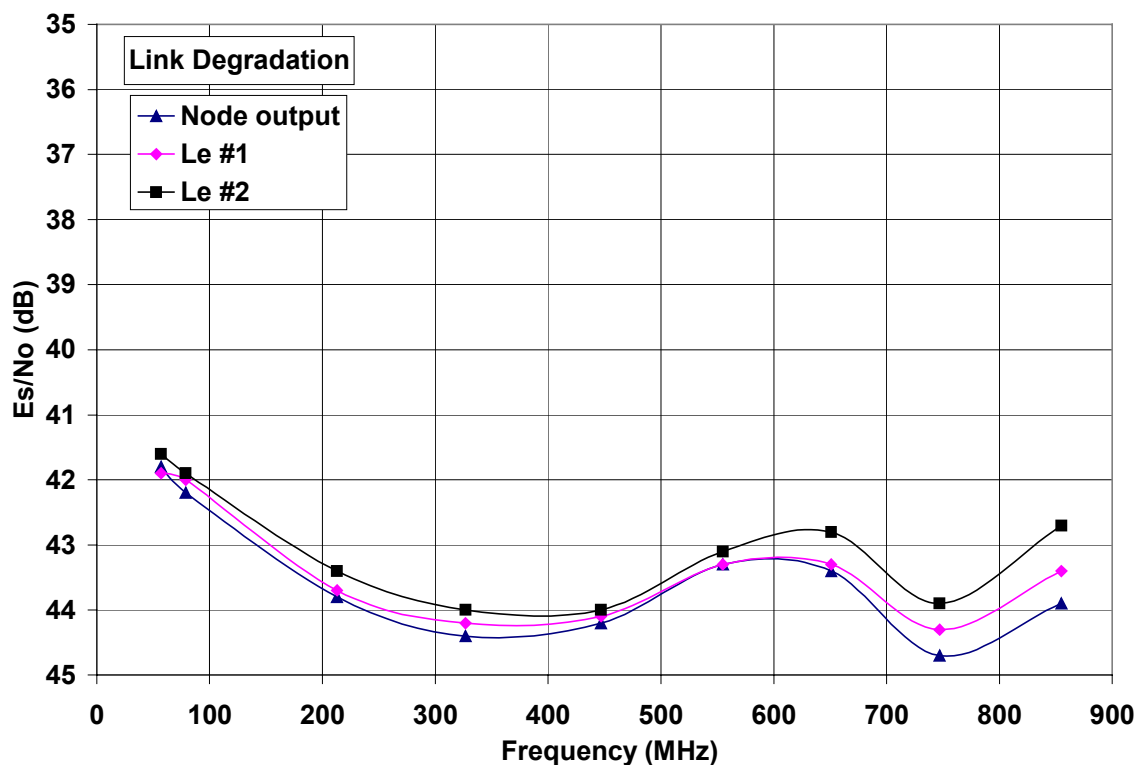


Figure 4 – Es/No versus Frequency for Distribution Depth

Table 4 – Distribution System Depth Effect on MER

ALM-11 Preset - Video Mode							
EIA Channel	Center Frequency (MHz)	SG-2000 MER (dB)		Line Extender 1 MER (dB)		Line Extender 2 MER (dB)	
		64 QAM	256 QAM	64 QAM	256 QAM	64 QAM	256 QAM
2	57	34.7	33.6	34.6	33.4	34.4	33.3
5	79	34.6	33.5	34.6	33.6	34.6	33.6
13	213	34.3	33.6	34.8	33.8	34.8	33.7
41	327	34.8	33.8	34.9	33.7	34.7	33.4
61	447	35.1	33.9	34.9	33.7	34.9	33.7
79	555	34.9	33.8	34.9	33.8	34.8	33.8
100	651	34.9	33.7	35.0	33.8	34.8	33.9
116	747	35.0	33.9	35.0	34.0	35.2	33.9
134	855	35.0	33.9	34.8	33.6	34.9	33.6

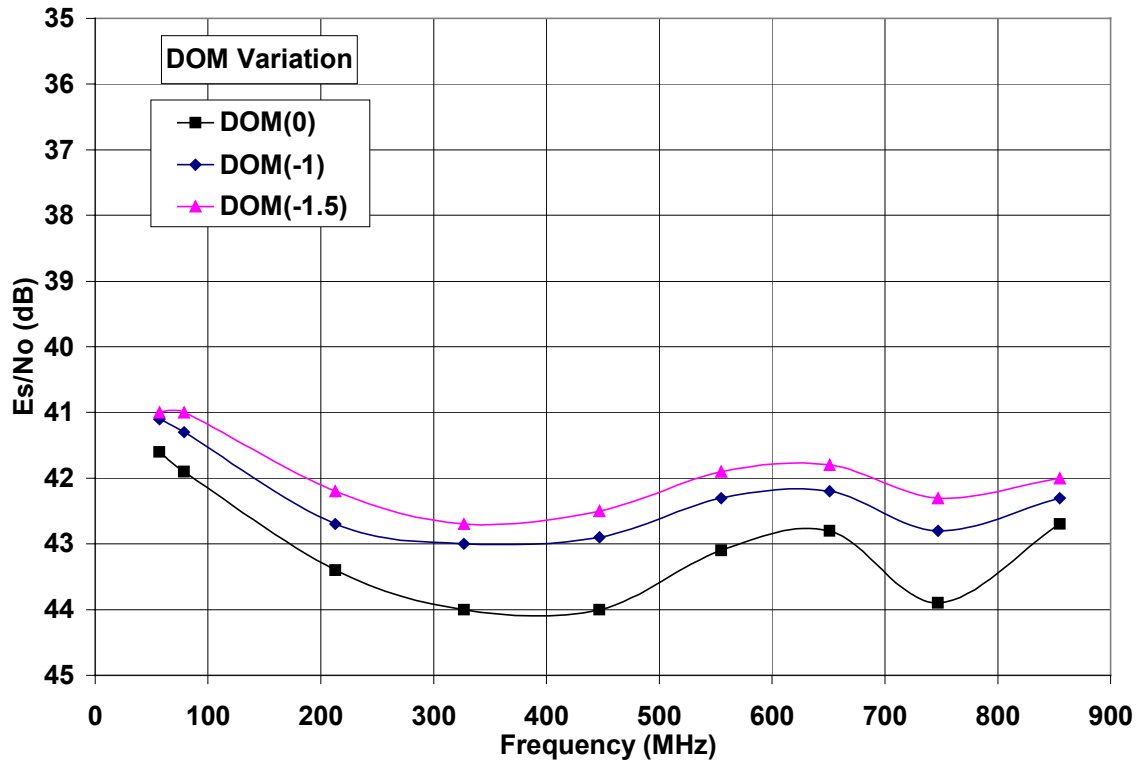


Figure 5 – 64 QAM Es/No versus Frequency for DOM Variations

Table 5 – DOM Variations on 64 QAM MER

64 QAM							
EIA Channel	Center Frequency (MHz)	DOM (dB) / LDL (dBm)					
		0 / 5.75		-1 / 4.75		-1.5 / 4.25	
		MER (dB)	Pre-FEC BER	MER (dB)	Pre-FEC BER	MER (dB)	Pre-FEC BER
2	57	34.4	0	34.4	0	34.2	0
5	79	34.6	0	34.5	0	34.4	0
13	213	34.8	0	34.7	0	34.3	0
41	327	34.7	0	34.6	0	34.7	0
61	447	34.9	0	34.8	0	34.8	0
79	555	34.8	0	34.7	0	34.6	0
100	651	34.8	0	34.7	0	34.7	0
116	747	35.2	0	34.7	0	34.8	0
134	855	34.9	0	34.7	0	34.7	0

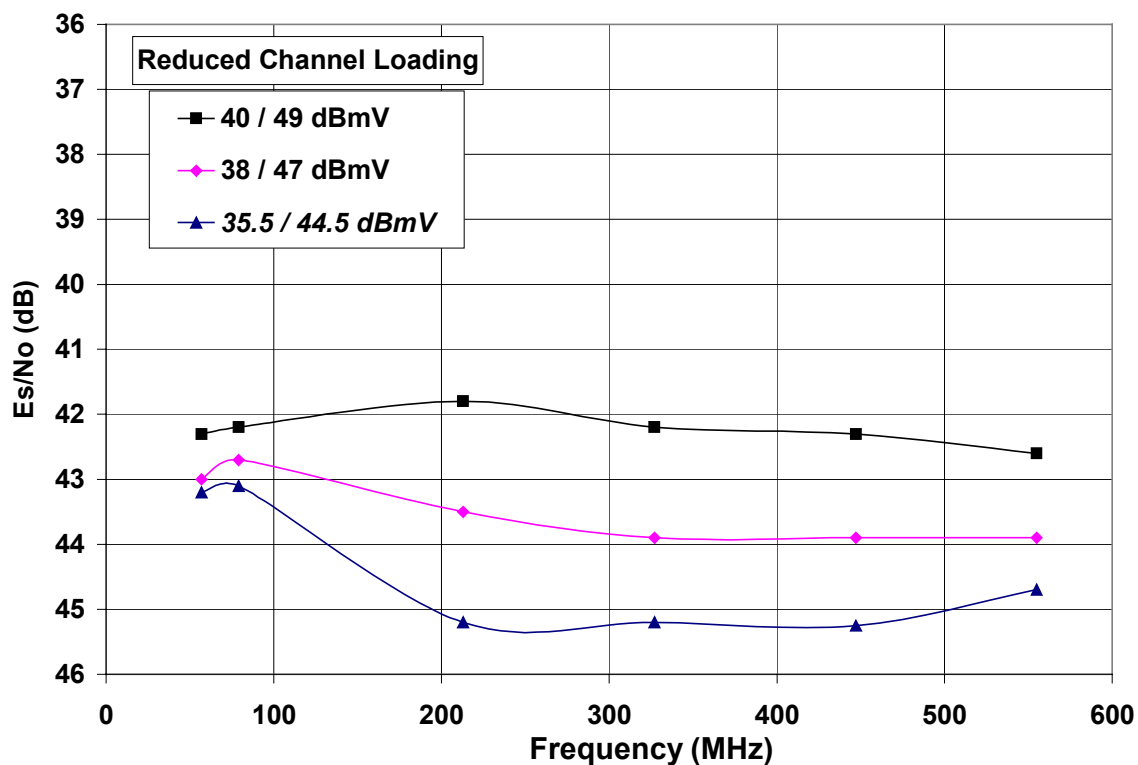


Figure 6 – 64 QAM Es/No versus Frequency for Reduced Channel Loading

Table 6 – 64 QAM MER for Reduce Loading on the Distribution System

Laser Preset Video Mode - Reduced Loading												
Center Frequency (MHz)	SG-2000 RF Output Level 55 / 603 MHz (dBmV)											
	40 / 49				38 / 47				35.5 / 44.5			
	256 QAM		64 QAM		256 QAM		64 QAM		256 QAM		64 QAM	
	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER	MER (dB)	Pre-RS BER
57	33.4	0	34.6	0	33.5	0	34.7	0	33.5	0	34.6	0
79	33.6	0	34.5	0	33.7	0	34.8	0	33.8	0	34.7	0
213	33.4	0	34.1	0	33.6	0	34.7	0	33.7	0	35.0	0
327	33.2	0	34.5	0	33.4	0	34.9	0	33.4	0	34.9	0
447	33.5	0	34.8	0	33.8	0	35.0	0	33.8	0	35.1	0
555	33.8	3.5E-08	34.8	0	33.9	0	34.9	0	33.9	4.8E-09	35.2	0

Table 7 – 256 QAM Long Term Bit Error Testing

256 QAM - Long Term BER				
LDL (dBm)	5.75	5.75	4.25	5.75
DOM (dB)	0	0	-1.5	0
Loading	55 – 855 MHz	55 - 855 MHz	55 - 855 MHz	55 - 603 MHz
LE #2 Output Level (dBmV)	32.5 / 43.5	35 / 46	32.5 / 43.5	40 / 49
EIA Channel	134	100	116	79
MER (dB)	33.6	33.64	33.8	33.73
Pre-RS BER	1.80E-09	2.68E-09	3.00E-09	3.64E-09
Post-RS BER	0	0	0	0
Error Seconds	0	0	0	0
BER	0	0	0	0

NETWORK ADDRESS TRANSLATION IN HOME GATEWAYS

Jed Johnson

Art Harvey

Motorola, Inc.

Abstract

Network Address Translation (NAT) has become a common feature in Home Gateways because it reduces the number of IP addresses a service provider needs to manage, it bounds the domain of the network and it provides a modicum of security for the home owner. NAT however does introduce problems because it breaks the end-to-end addressing assumption built in to many applications. In addition, the applications that have the most difficulty are generally the applications that deliver advanced services like IP telephony and streaming media.

NAT also makes network management difficult because the devices behind NAT cannot be addressed directly. If a service provider wants to be able to diagnose a network problem through NAT, many standard tools and procedures will not work.

Because of these problems, NAT in many circles has been equated to "a bad thing" that must be eliminated. This paper takes the position that NAT cannot be eliminated from all, and some might say most, home networks so we should learn to deal with it. This paper specifically looks at how end-to-end management and advanced services can be delivered with NAT in place in the home network.

This paper will review how NAT causes problems and then go on to show how extensions or work-arounds to NAT can recover the end-to-end addressing assumption that applications require to work properly.

INTRODUCTION

NAT has become prevalent in home networks so it needs to be discussed

We are not NAT fanatics but believe it can serve a useful purpose in home gateways

This paper will provide some background on NAT including a brief overview of how NAT works. The paper will discuss some advantages that NAT has and a section on debunking myths about NAT is included. Finally the paper covers some system level descriptions of how NAT can be used to solve unique home networking problems.

NAT OVERVIEW

The most concise definition of NAT is, "Network Address Translation is a method by which IP addresses are mapped from one realm to another, in an attempt to provide transparent routing to hosts." (RFC 2663) This definition is typically implemented as a gateway device that connects a private address space, such as that of a homeowner or business, to the public Internet address space through an Internet Service Provider (ISP). The NAT gateway replaces a private network address with a public one in packets sent from a system in the private network to the public network, and performs the inverse replacement for packets flowing in the reverse direction. This mapping of addresses between addressing realms is called "transparent routing".

RFC 2663, *IP Network Address Translator (NAT) Terminology and Considerations*, presents an overview of the variants of NAT

and the standard terminology. It also describes the characteristics of NAT, typical usage, operational characteristics and limitations. We only briefly describe some of these here and then mainly in the context of home networking.

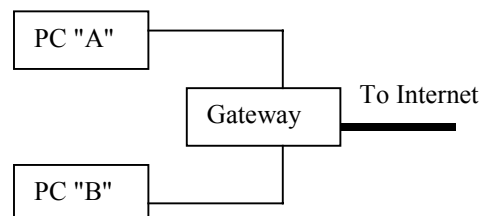
There are many reasons for using NAT. A common one for large organizations is to avoid the problem of changing the network address of every system in the private corporate network if there is a change of the set of addresses provided by their ISPs. For the home network, the prime motivation for NAT is to share the single public network address provided by the ISP among multiple systems in the home so that all the devices in the home have Internet access.

There is no single method or standard for NAT and the variations are many. We present conceptual overviews of three of the main variants relevant to home networking. We skip many of the details, but more complete descriptions and other variants can be found in the appropriate RFCs and IDs. The methods we describe are known as:

- NAT with dynamic address assignment, see RFC 3022, *Traditional IP Network Address Translator (Traditional NAT)*
- NAPT - Network Address Port Translation, see RFC 3022, *Traditional IP Network Address Translator (Traditional NAT)*
- RSAP-IP - Realm Specific Address and Port IP, see draft-ietf-nat-rsip-framework-05, *Realm Specific IP: Framework*; and draft-ietf-nat-rsip-protocol-07, *Realm Specific IP: Protocol Specification*

Consider the case of a home with two PCs, a Home Gateway, and a connection to the Internet as shown in the following figure. Further suppose the ISP providing connection

to the Internet assigns a single IP address to the Home Gateway.



While data can be exchanged within the home using privately assigned addresses (RFC1918, *Address Allocation for Private Internets*), these addresses may be duplicates of those used in some other private realm and cannot be routed outside the home LAN. Suppose the Home Gateway supports NAT with dynamic address assignment. If PC "A" sends a packet to a remote system in the public network that has a globally assigned (unique) address, then the NAT gateway binds (typically stores in a mapping table) the private IP address of "A" to the IP address assigned by the ISP. Next, it replaces "A's" private IP address in the packet with the IP address assigned by the ISP. It then forwards the packet to the public network. If the gateway subsequently receives traffic from the remote system, it performs the inverse mapping (replaces the destination IP address with the address of "A") and forwards the packet on the private network. Based on some heuristic, (for example, receipt of a TCP FIN message indicating the connection is terminated and allowing some time for retransmission of lost packets) it unbinds (removes from the mapping table) the address of "A" from the IP address assigned by the ISP. "B" can now go through the same steps to communicate outside the private network.

Note that while "A" is using the single ISP assigned address, "B" cannot send or receive data from the public network and vice versa.

NAPT, Network Port Address Translation, enables multiple concurrent systems to communicate with remote systems across the public network. It does this by replacing not only the IP address of packets crossing the border between the private and public networks but also the TCP or UDP port address. The following example demonstrates how this typically works.

Suppose "A" has a locally assigned IP address of **IP-A**, and "B" has locally assigned IP address of **IP-B**. Denote the single, shared IP address assigned by the ISP as **IP-External**.

Consider the case where both hosts, "A" and "B", establish a connection with the same remote host, that has a global IP address of **IP-Rem**, and to the same application on a well-known TCP port denoted **TCP-ServerPort**. Also, assume both "A" and "B" choose the same source TCP port number **TCP-ClientPort**. [The algorithm operates in the same fashion when "A" and "B" choose different source port numbers and talk to different systems.] To initiate the connections, "A" and "B" send packets that include the following IP and TCP header information.

"A" sends

Destination IP address = **IP-Rem**,
Source IP address = **IP-A**,
Protocol = TCP,
Destination Port = **TCP-ServerPort**,
Source Port = **TCP-ClientPort**,
TCP message type = SYN

"B" sends

Destination IP address = **IP-Rem**,
Source IP address = **IP-B**,
Protocol = TCP,
Destination Port = **TCP-ServerPort**,
Source Port = **TCP-ClientPort**,
TCP message type = SYN

The gateway NAPT receives these packets and notices the destination IP address in each packet is in the public address space and the

protocol is TCP. The SYN indicates this is a new connection. These packets cause the NAT to create a mapping table entry that will last until it sees FIN messages that terminate the TCP connections or by some other heuristic. The NAT uses this table to change the outgoing packet headers above to:

Packet sent by "A" is translated to

Destination IP address = **IP-Rem**,
Source IP address = **IP-External**, (*the external address to be shared*)
Protocol = TCP,
Destination Port = **TCP-ServerPort**,
Source Port = **TCP-A1** (*a different source port*)

Packet sent by "B" is translated to

Destination IP address = **IP-Rem**,
Source IP address = **IP-External**, (*the external address to shared*)
Protocol = TCP,
Destination Port = **TCP-ServerPort**,
Source Port = **TCP-B1** (*a different source port*)

The modified packets are then forwarded to the public network. To the remote system, these appear to be two different connection requests from a single host with IP address **IP-External**. Packets received by the NAT device from the remote system undergo the analogous inverse translation, and are sent to the private network.

The nice thing about NAT is that it performs transparent routing. The PCs "A" and "B", and the remote system have no idea that the gateway is modifying the addresses.

Unfortunately, this does not always work as desired. For example, the FTP application has messages ("PORT" and the "PASV" response) in which it includes the local system IP address in its data. Since a private address is of no use to a system outside that private network, a NAT ALG (Application Layer Gateway) must do the same type mapping of these IP addresses inside the application data as it does to the packet headers. This can get

tricky because these addresses are encoded in ASCII so that changing the address may also change the size of the packet. That means the ALG must also modify checksums, the TCP fields, and maintain state for TCP sequence numbers and acknowledgments. Throw IP fragmentation into the mix and it becomes apparent that maintaining transparency is not trivial. Furthermore, IPsec transport mode, both AH and ESP, include an integrity check over the entire payload including the TCP and UDP checksum. Modifying headers in protected packets will cause the receiving IPsec to discard the packet as having failed the integrity check. While these problems may seem insurmountable, they can all be addressed as described in subsequent sections.

Partially relaxing the transparency constraint eliminates many of these problems and is one of the motivations behind RSAP-IP (Realm Specific Address and Port IP). With RSAP-IP, address translation remains transparent to the application, but the network stack at the end systems are aware of the address mapping. Here is one possible implementation.

A system in the private network, say "A", queries an RSAP-IP server in the Home Gateway asking for an IP address and port number. The Home Gateway establishes the binding between the private IP address plus port of "A" and the external IP address plus port just as it does for NAPT. The Home Gateway responds to the query from "A" by returning this binding. Now when "A" sends a packet to a system outside the private network, "A" uses the external IP address and port numbers in the packet header. Also, if an application (such as FTP) asks for an address, the local stack can return an external address. Packets sent from "A" might be tunneled (encapsulated in another IP header) to the Home Gateway for decapsulation and transmission on the public network. They could also be tunneled from "A" directly to a remote system. They may even be sent just as

they are since the Home Gateway knows the address binding and can know how to route the packets. Other alternatives are possible as well.

ADVANTAGES OF NAT IN A HOME GATEWAY

Firewall by nature

While many networking purists cringe at the thought, NAT in a broad sense can be regarded as a firewall and is sold as such in some Home Gateway products. Because most NAT implementations in Home Gateways only open up ports based on traffic that is initiated in the home, the only ports that are open are to support applications that reside in the home. Unsolicited traffic to any another port is dropped. The only traffic that gets through is for the ports associated with applications in the home and only while those applications are running.

There are variations of NAT that open up ports for UDP traffic and for servers in the home. These features obviously reduce the effectiveness of NAT as a firewall and need to be used with this fact in mind.

Natural demarcation point of ISP

The ISP providing Internet access to the home might not want to get involved with supporting the home network. If a homeowner has multiple PCs and wants each PC on the network then without NAT each PC needs a separate address and in all likelihood these PCs are located in different rooms in the house. This means that from an IP standpoint the ISP can be expected to managed connectivity through the home network to each PC.

The problem for the ISP is the lack of physical access to the home network and local of configuration control.

By using NAT in the Home Gateway the ISP can terminate management of IP connectivity in the Home Gateway and leave the

management of connectivity in the home to the homeowner.

Allows device provisioning in the home without direct involvement of ISP

Without NAT in the Home Gateway each device that is added to the network needs to be provisioned by the ISP. There are automated provisioned systems to reduce the manual effort for this process but resources like IP addresses are required in greater abundance.

With NAT in the Home Gateway IP management can be terminated in the Home Gateway as discussed in the previous section. This allows device provisioning up to the IP layer to occur in the home under the control of the Homeowner and transparent to the ISP.

It is possible, by the way, for the ISP to limit the number of devices that access the Internet from the home in this scenario. This means the ISP does not lose out on implementing a potential business model by not participating in network provisioning.

It also needs to be pointed out that service provisioning at the application layer is another subject and it can be independent of network provisioning.

Eases migration to IPng.

NAT can translate between different network layer protocols. Any migration from IPv4 to another address format will require this for existing systems to continue working. RFC 2766 *Network Address Translation - Protocol Translation (NAT-PT)*

Since proxies, firewall, other gateways will be at the boundary, NAT is easily included with little additional overhead since they are either terminating the connections, or already inspecting and modifying headers and application data. This also relates to myths that NAT is too slow.

DEBUNKING THE MYTHS

Myth 1: Devices behind NAT cannot be managed

It is true that devices behind NAT cannot be managed using traditional methods that depend on public IP addresses. However, there are mechanisms that get around this issue.

There are RFCs and Internet Drafts that describe how to managed devices with SNMP through NAT. It can be complicated in some cases but it can be achieved.

There are also mechanisms that use SNMP Proxy Agents, discussed later in this paper, that also provide the ability to manage devices behind NAT. The SNMP commands are terminated on the public side of NAT and SNMP is not required in the end devices. This can be an advantage for some devices.

Myth 2: Can't support multiple ISPs with NAT

It has been claimed that mechanisms like source routing that are used to support multiple ISPs don't work with NAT. While we have not investigated every possible means of supporting multiple ISPs we have shown that source routing can be used with NAT to support multiple ISPs. A section later in this paper goes into more detail.

Myth 3: NAT breaks end to end security

Security mechanisms may be applied to any layer of the communication architecture. All mechanisms above the network layer, Transport Layer Security (TLS) for example, do not interfere with the NAT address mapping in the network layer. However, TLS does preclude ALGs from modifying the IP address information that may be present in some application protocols. As described later, this is not a problem when using RSAP-IP.

IPsec secures packets at the network layer. IPsec security includes protecting the integrity

of the parts of the IP header not modified by routers and all data contained in the IP packet. Integrity protection guarantees that any modification of the data will be detected and the packet will be discarded by the receiver. As the main function of NAT is to modify the IP (and possibly TCP or UDP) addresses and header, IPsec and NAT directly conflict.

There are two common solutions to the problem. If the NAT gateway is at border between a trusted, private network, such as the home, and the untrusted, public network, then tunnel-mode IPsec alleviates all difficulties (RFC2709). Data is encapsulated and protected between the NAT gateway and the system across the public network. Packets sent from the home are first subjected to NAT, then IPsec protections are applied. Incoming packets for the home are received by the gateway, the IPsec protection is checked and removed, the address is translated, and lastly the packet is forwarded on the home network. If the home network is subjected to threats that demand safeguards, then the gateway can establish an IPsec tunnel to the system in the home so that all packets are protected there as well.

Where protection must extend from the system in the home all the way to the remote system with no intermediaries, then the RSAP-IP is an alternative to gateway tunneling (draft-ietf-nat-rsip-ipsec-04, *RSIP Support for End-to-end IPsec*). As described in the overview, the system in the private network acquires an external IP address (and TCP or UDP port number if needed) from the RSIP server. The system forms packets in the normal way using this external address. Any IPsec protections are applied by the system before transmitting the packet. Since the packet already contains the externally routeable IP address, the gateway no longer modifies the IP or TCP/UDP headers, and IPsec operates end-to-end.

Myth 4: NAT means you can't deploy advanced services

We have heard that service providers resist or reject using NAT in Home Gateways because advanced service cannot be deployed using NAT. Typically these services require multiple ports, some using UDP and typically these services need to support asynchronous traffic into the home.

It is true that basic NAT cannot support these applications but there are two well known approaches to supporting applications like these that have been able to support every application we have encountered.

One approach is the Application Level Gateway for ALG. This is an application aware algorithm that runs within NAT to provide assistance in address translation and port binding. This algorithm knows the content of the messages and provides the translation within the messages as required.

Another approach is to put part of the application into the Home Gateway. A section later in this paper shows an example using SNMP. Another popular example to support IP telephony using protocols like H.323 by putting part of the telephony application in the gateway.

We will grant that these approaches are different and by no means traditional but they preserve the advantages of NAT without breaking the application.

Myth 5: NAT is too slow.

There is no doubt that address translation takes time and reduces throughput, however, if Home Gateways become the residence of a firewall to protect the Home Network from the Internet then we claim the overhead to examine IP packets has already been put in place and the incremental time needed for address translation is negligible.

The assumption of course is that IP packet headers are being processed for another reason

independent of NAT. If this is not the case then NAT will reduce throughput especially when compared to a level 2 switch. NAT in IP routers can also be shown to reduce throughput but to a lesser degree.

SYSTEM LEVEL NAT EXAMPLES

Managing devices behind NAT

Adding capabilities around NAT can create a full set of management capabilities for a range of device types in the private address space behind NAT. One approach is to deploy an SNMP Proxy Agent that as an application has access to both the private and public address spaces on each side of NAT. The figure below shows how this could be designed.

In this approach an SNMP Proxy Agent opens up two network interfaces; one for connecting to the public address space and the other for connecting to the private address space.

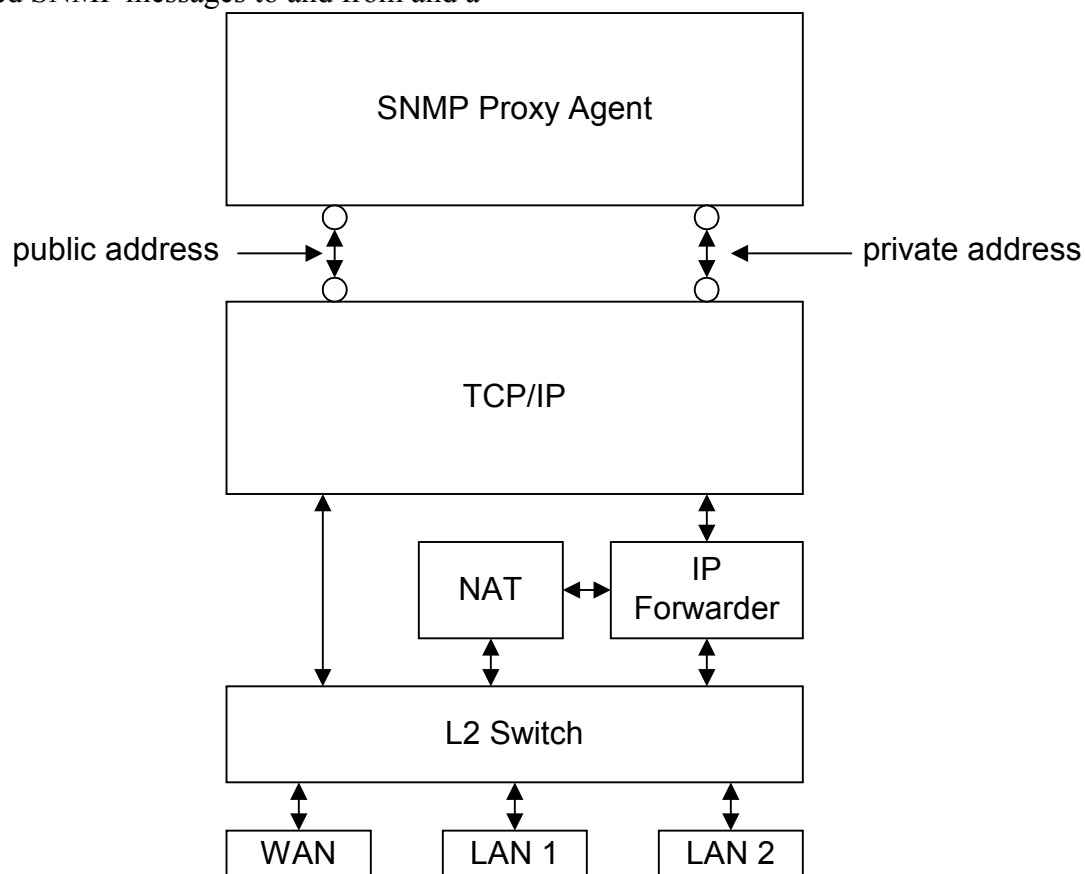
The port on the public side is used to send and received SNMP messages to and from a

management system also in the public address space, presumably located at the MSO head end or network control center.

The port on the private side is used to send messages in an arbitrary format to the appropriate devices or objects in the private address space.

In this configuration NAT is not part of the data flow. NAT is in the system to act as a quasi-transparent address translator for end-to-end applications. In the SNMP Proxy Agent case, the TCP connections are terminated and appropriate addresses that do not need translating are used.

The SNMP Proxy Agent can use an approach that provides a separate Object Identifier (OID) for each managed object or device and thereby give the look to the management system that each object has its own SNMP agent. A private MIB is created for each Object class.

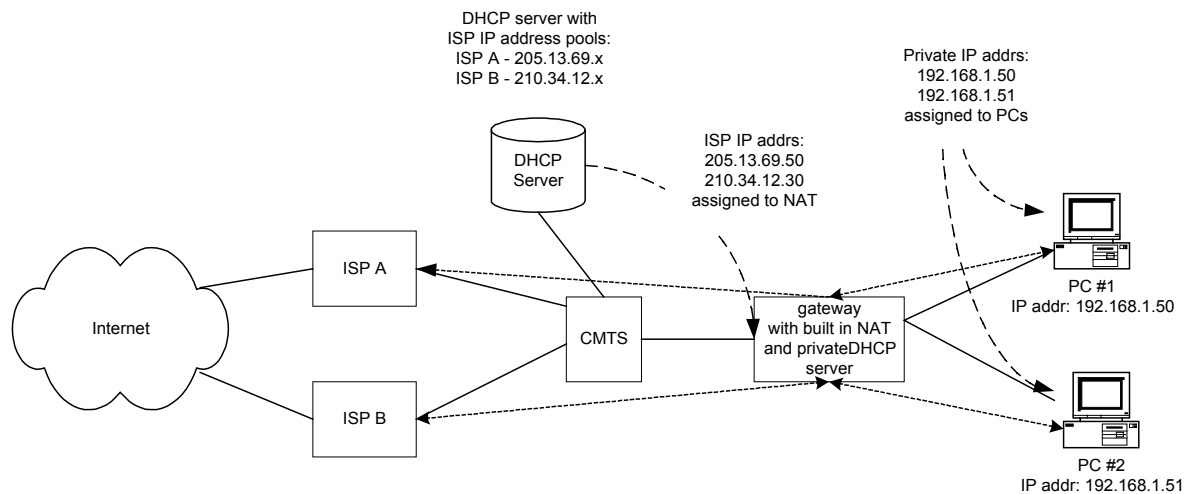


Multiple ISPs with NAT

Overview

In a Home Gateway with NAT, the Home Gateway translates and forwards all IP traffic to and from the CMTS and the customer premise equipment (CPE). The CPE is a home owner PC. The following list describes the provisioning process when multiple ISPs are used:

- 1) At startup, the PC transmits a DHCP request for a private IP address from the Home Gateway's DHCP server.
- 2) The PC binds its MAC address with that private IP address assigned to it.
- 3) Since no ISPs have been provisioned, the only web site the end user can get to is a web site to provision an ISP. The end user provisions an ISP using this web page.
- 4) After the ISP is provisioned, the Home Gateway is forced to have NAT get a new IP address. This can be done having the provisioning server send a command to the Home Gateway forcing it to release and renew the IP address.
- 5) When NAT attempts to get a new IP address, the MAC address associated with NAT is assigned an IP address for the provisioned ISP. This is because the DHCP server associated with the provisioned ISP has been configured with the MAC address.
- 6) The NAT in the Home Gateway now translates all IP traffic from that PC to the selected ISP IP address.
- 7) At this point all the PCs in the house would use this ISP.
- 8) If the End User would like to provision another ISP the user will open a management window to the Home Gateway and request to provision another ISP.
- 9) The NAT function would use the default IP address for the provisioning server. The PC in use by the end user would at this point only be able to go to the provisioning web page.
- 10) The end user would provision another ISP and NAT would be forced to renew the IP address for the 2nd ISP.
- 11) At this point NAT would have two IP addresses. All the PCs except the one used to provision the new ISP would be connected to the old ISP. The PC used to configure the new ISP would be connected to the new ISP. A management window to the Home Gateway would be available for the end user to move PCs between ISPs.
- 12) Additional ISPs can be configured in the same way.



Discussion

IP addressing on the Home Network is handled using a local DHCP server. The address space is private and NAT is used in the Home Gateway. All PCs in the home are on the same subnet and bridging between networks in the home is performed by the Home Gateway.

Before any ISPs are provisioned NAT has one public IP address and it can only be used to access the provisioning server. The end user can use a browser on any PC to access the ISP provisioning server and provision ISP.

Once an ISP is provisioned, the Home Gateway needs to get a new IP address. The preferred method of doing this is to have the provisioning server send an SNMP command to the Home Gateway that would force the Home Gateway to release the current address and then requested a new one. When the DHCP server renews the address, it will provide one for the provisioned ISP. The proper address is obtained because the provisioning process configured the DHCP server with the MAC address from the NAT function.

At this point all the PCs in the home access the Internet through the same ISP.

If another ISP needs to be configured for the home the user opens up a management window to the Home Gateway and requests a new ISP. NAT can use a single MAC address and share it across multiple IP addresses or NAT can allocate another MAC address. The choice here depends on how the DHCP server works and whether it can handle one option or the other. The preferred is to minimize the use of MAC addresses that NAT needs.

Either way, NAT binds the private address of the PC that made the request to a public address. At this point, that PC can only access the provisioning server.

After the new ISP is provisioned NAT gets another public IP address and now all the PCs in the home except the one that provisioned the new ISP are connected to the old ISP. The PC that was used to provision the new ISP is connected to the new ISP.

A management window to the Home Gateway can be used by the end user to configure which PCs in the home are connected to which of the ISPs.

Additional ISPs can be configured in the same way.

RFCs

IP Network Address Translator (NAT)
Terminology and Considerations (RFC 2663)

DNS extensions to Network Address Translators (DNS_ALG) (RFC 2694)

Security Model with Tunnel-mode IPsec for NAT Domains (RFC 2709)

An SNMP Application Level Gateway for Payload Address Translation (RFC 2962)

Traditional IP Network Address Translator (Traditional NAT) (RFC 3022)

Protocol Complications with the IP Network Address Translator (NAT) (RFC 3027)

Network Address Translation - Protocol Translation (NAT-PT) (RFC 2766)

ABOUT THE AUTHORS

Jed Johnson

Jed.Johnson@motorola.com

Jed Johnson is the Director of Engineering for Home Networking products in Motorola's Broadband Communications Sector (formerly General Instrument). Jed has been with Motorola for 17 years and has held management and technical positions in various organizations including Home Networking, Streaming Media, Cable Data Products, Transmission Products and Network Management. Previously Jed worked at The Foxboro Company for 7 years as both a Software Engineer and Hardware Engineer. He began his career at Data General.

Jed has a BET degree from Northeastern University and an MBA from the University of Massachusetts at Dartmouth.

Art Harvey

Arthur.Harvey@motorola.com

Arthur Harvey is the Director of the Broadband Networks Research Laboratory within Motorola Labs and has been with Motorola for two years. He worked previously at the Open Software Foundation as Director of the Architecture Group. Prior to that, Arthur was a Network Architect at

Digital Equipment Corporation (now Compaq).

Arthur has Ph.D. in Biophysics from the University of Virginia .

NETWORK TRAFFIC MODELING AND PLANNING

Performance Verification of Cable Networks

Hardev Soor, Ixia

Abstract

The dramatic growth in the volume of traffic over cable modems has increased the demand for traffic modeling and planning mechanisms for verifying cable network performance via traffic generators and analysis platforms. Packet generators of today offer highly sophisticated utilities to generate custom traffic mixes for Layer 2 and Layer 3 devices in an effort to simulate real-world Internet traffic patterns. In addition, such analysis systems assist cable network engineers in developing and deploying cable systems in the networks by measuring their performance and gathering metrics such as throughput, packet loss, and latency.

Within the past few years, Cable service operators have begun to integrate data services into their existing cable infrastructure. This is to fulfill the customer's desire for high-speed Internet access at a reasonable price. Also, this added service creates additional revenue for the MSO.

As the result of data services over cable, three major entities have emerged. They are the Small Office/ Home Office, Voice over Cable and Digital Video.

The Small Office / Home Office (SOHO) has benefited the most as a result of Data over Cable. Many companies have created network devices for the SOHO market.

Digital Cable, another result of Data over Cable improves the viewing quality of Television. Broadcasting downstream is a thing of the past. Multicasting provides an efficient way of delivering quality video to the subscriber.

Finally, voice over cable is the future of Data over cable, also known as PacketCable. Using this technology, a cable subscriber is able to make voice calls from a cable modem to anyone in the world.

Cable and DOCSIS

As Cable Modems are in its adolescent stages, CableLabs is conducting conformance tests by certifying Cable Modems and qualifying Cable Modem Termination Systems by verifying the basic functionality of the devices.

For the MSO, since there are so many Cable modems and CMTS's out there, they need to verify performance at the device level as well as the system level. It's up to the MSO to evaluate the performance of each cable device to find out which of the devices meets their requirements. It is also a requirement of the MSO to find out how these devices behave in a system environment. A typical cable company handles around 5000 to 30,000 subscribers.

Testing Requirements of Cable Device and Systems

As the development of cable networking devices gets more complex, performance analyzers must meet the conformance requirements of these manufacturers. In addition, as the MSO's implement the devices into their infrastructure, Performance Analyzers must meet the needs of this group as well. The Performance analyzer must take a role in assisting the MSO to provide as much information as possible on the following:

- * Test the performance of the Cable Networking Device.

- * Test the performance and identify any interoperability issues of any of the Cable devices that participates in the Cable System.
- * Simulate the possible real world scenarios of the Cable system and assist in prestaging the system before deployment.
- * Monitor and provide valuable information of the Cable network
- * Identify any major problems after the system has been deployed.

Device Testing

In testing the performance of Cable Networking Devices, the following performance metrics of Throughput, Packet Loss and Latency should be used.

The device that consistently performs well in all of the three above categories should be considered. Device testing is very tricky because the performance of a Cable Modem is dependent upon the CMTS and vice versa. Therefore, the performance analysis of all combinations of different CMTS's and Cable Modems should be considered.

System Testing

Testing a typical system involves a CMTS with many modems. Questions arise when testing these systems.

For example:

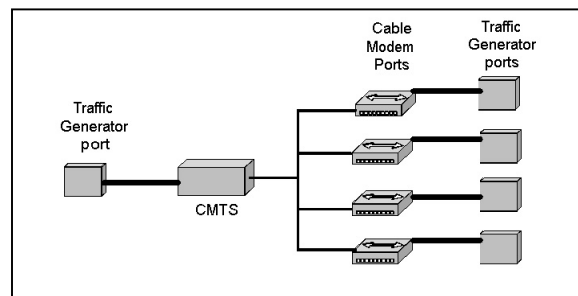
What is the maximum number of cable modems that meets the performance requirement of the CMTS?

What is the total aggregate throughput of a Cable System?

Is the system able to handle the basic requirements of the MSO like Throughput, Packet Loss, and Latency?

The Customer Premises Equipment (CPE) is another variable introduced into the system. The testing device must be able to simulate 1 or more CPE's per modem and on the CMTS side.

In the diagram below, the Traffic Generator/ Performance Analyzer is able to represent 1 or more CPE's on the same physical port.



The system that performs well to these questions should be considered.

Simulation of a Cable Network

Before deployment of a Cable Network or any Network for that matter, the equipment and personnel involved, must be prepared to identify any major flaws, improvements made to the system and expansion. A testing device must be able to simulate these scenarios in order to reduce downtime. In addition, it must provide information to the network administrator in order to make any improvements to that network. And finally, increase productivity of the network. The testing device should provide a solution to the following simulations:

- * Quality of Service (QoS)
- * Multicast

- * Voice over IP
- * Traffic shaping and modeling

Network Monitoring

To provide added value, the Test equipment should not only be used as evaluating Cable devices and systems, but also be included in the “nervous system” of the Cable Network. It should be able to warn the administrator of any security and fault issues in the network. Also, it should be able to provide performance analysis values of the whole network.

The performance analyzer should be able to filter packets that are meaningful as well as warn the administrator of the occurrence of an event. Examples include SYN attacks, severe errors, or link failure.

Traffic Modeling

Of all the testing devices and systems are being utilized today, most cable companies are most interested in the topic of traffic shaping and modeling. Traffic generators / performance analyzers today are able to simulate traffic based on a real cable network. To simulate the types of Internet traffic, one must find out what types of traffic is being passed through the network. Some types of traffic include the following:

- * DHCP traffic-- as Cable Modems go through their registration process, they communicate via DHCP to the DHCP and TFTP servers to get their configuration files and IP addresses. CPE's also use the DHCP protocol to get IP addresses.
- * Web traffic-- as users make HTTP requests to different web servers throughout the Internet. Web traffic is based on TCP sessions.

- * Domain Name Service (DNS)-- is used to resolve a name with IP addresses. As more and more users access web pages, the web addresses are resolved by the DNS sever.

- * Email traffic-- Most people have joined the bandwagon and communicate via e-mail in today's demand for high-speed communication. Different types of mail protocols generated are POP, SMTP, and IMAP to name a few.

- * Real-time Transport Protocol (RTP)-- As digital cable and voice are integrated into the cable network, the information will be encapsulated into this protocol.

- * File Transfer Protocol (FTP) -- Subscribers will definitely download or upload files between each other. Examples include download of technical documents, video and music files to name a few.

Once the matter has been researched, to measure the performance of the aggregate network, different types of proportions of traffic need to be generated in order to accurately simulate a typical cable network. In other words, the MSO needs to figure out the percentage of the specific types of traffic that must be generated.

The performance analyzer is used to generate custom traffic that can closely mimic real-world traffic, where the users send and receive IP traffic from their web pages. To simulate these web pages, the device needs to generate TCP, UDP, and HTTP traffic. The important characteristics of such traffic generation are:

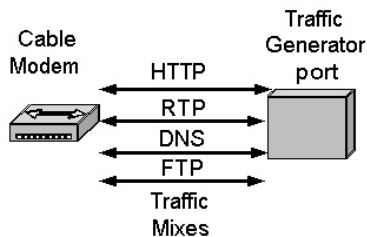
- * Variable packet sizes
- * Variable allocated bandwidth

- * Priority (IP Type of Service)
- * Bursty and fixed traffic

The testing procedures that simulate such characteristics can be complicated in nature.

Variable packet sizes are necessary because different types of traffic require different packet sizes. For example, TCP connections require small packets since the information exchanged is done by the TCP header and the payload is not necessary.

Voice packets require small size packets to maintain low latency and jitter. UDP traffic requires large packets because of the need to transfer as much information possible with less overhead. DHCP packets require medium-sized packets since the nature of transferring all the DHCP options.

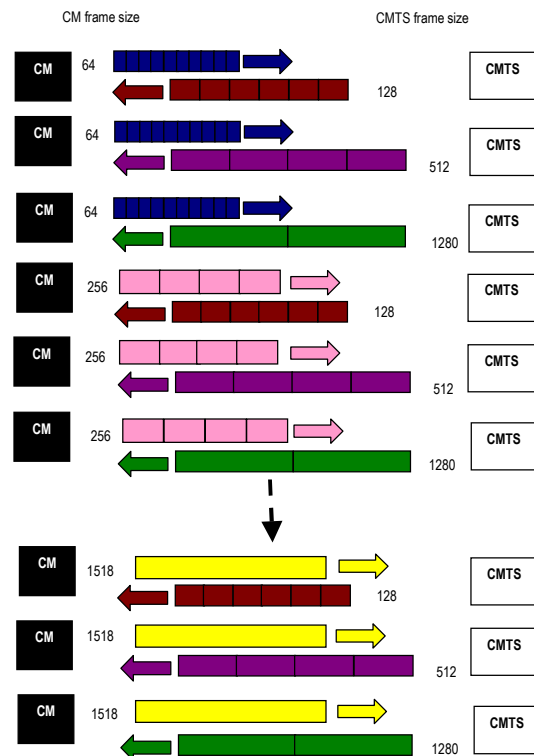


Variable allocated bandwidth is taken into consideration when the ratio of traffic is different. The ratio of traffic also changes over time.

It is also necessary to generate bursty and fixed traffic. For example with voice packets, a fixed constant rate is needed to maintain constant latency and jitter. If the voice traffic is not constant, the voice quality decreases. Also, since multiple users randomly send and receive traffic, bursty types of traffic can closely simulate this situation.

In order to simulate real-world traffic, the device must generate randomized types of traffic. So variables will be randomized, the packet size, the type of traffic represent both

the IP type (TCP/UDP) and their port numbers that represent HTTP, DNS, FTP etc.



In a Bidirectional test, frame sizes and the type of packets do not have to be the same.

<i>Frame Size</i>	<i>Protocol</i>	<i>Precedence</i>	<i>% Bandwidth</i>
80	TCP	3	20
64	RTP	1	10
570	DHCP	5	30
128	FTP	4	20
1518	UDP	8	20

Quality of Service

Quality of Service (QoS) is implemented in DOCSIS 1.1. With QoS in the cable modem world there are 2 terms associated with it. They are Service Flows and Classifiers. Different types of bandwidths are set on each Service flow to set priorities on

which types of packets are sent first. Each Service flow is based on one or more classifiers. The classifiers are based on different byte values at different offsets of the packet. Examples of classifiers are TCP/UDP ports, IP addresses, MAC addresses, and TOS Bits.

The analyzer used must be able to generate and analyze the traffic that matches these classifiers.

Conclusion

In conclusion, this article has discussed a variety of requirements necessary for a typical performance analyzer. The performance analyzer must be able to prepare and improve the quality of the Cable network in evaluating the network devices, simulate the environment in pre-deployment and assist in gathering performance metrics, as well as warn the administrator of any issues in post deployment

NEW CARRIER CLASS ARCHITECTURE BUSTS BANDWIDTH BLUES

Paul Connolly
Scientific-Atlanta, Inc.

Abstract

Network operators are on the burgeoning edge of providing interactive video, data and voice services. In offering these technologically advanced services, operators are in search of an enabling platform to deliver quality of service and reliability for interactive applications.

This paper explores optical network platforms and their ability to provide broadband operators with the technologies needed to take advantage of their unique position to deliver interactive video, data, and voice services, previously associated only with today's most advanced telecommunications systems. Discussions will offer new solutions allowing operators to deliver the bandwidth advantages of optical networking closer to subscribers' homes.

THE OPPORTUNITY

Broadband network operators have spent the past few years transitioning their networks from one-way analog entertainment to multimedia interactive networks. Voice and data services have been scaled to reasonable volume, and video has been expanded from analog to include digital -- and now interactive digital video services. Unlike other telecommunications network transitions, this transition has been fueled by a clear willingness to pay for new services by the users.

Now, however, the very success of these new offerings is leading to more demands by customers for bandwidth intensive services. The first generation of

technologies used to build these interactive networks will not be adequate as the bandwidth growth pushes these technologies beyond their breaking points.

A number of new technologies are fortunately becoming available to enable significant bandwidth capacity increase to be engineered in a cost effective manner.

CURRENT NETWORK LIMITATIONS

With approximately one percent of U.S. households currently using voice services on HFC networks, four percent using data, and less than one percent using video-on-demand, the cable industry is still in the very early stages of interactive deployment. However, it is clear that a number of impediments to high volume ramp currently exist. Beginning with voice services, the use of circuit switched technology requires a unique conversion device at each user's home, and coupled with network powering demands, has lead most network operators to forego the voice opportunity until voice can be integrated with the data network.

For data services, the industry has transitioned from proprietary systems to DOCSIS 1.0 and is currently preparing to move to DOCSIS 1.1-based networks. Cost of deployment has come down significantly via the standardization process, and with DOCSIS 1.1, capability will move beyond best efforts data to include managed data and voice services.

A key issue however is users future expectations on data rates. Current market pricing for CMTS functionality is typically in the \$500 per megabit forward and \$1,000

per megabit reverse capacity. For this reason, networks are typically engineered for less than 20 kb/s usage bit rate per home passed. The "bursty" nature of data, low penetration levels, and current usage rates, enable a relatively high data rate to be typically delivered. As more users turn to streamed data, whether for audio or especially video, average bit rates per user will climb dramatically. Current CMTS architectures will have great difficulty scaling to meet this demand. Furthermore, lack of ability to transport high numbers of multiplexed analog upstream paths has caused operators to deploy CMTS systems in hubs rather than headends. This results in strain on hub space as data rates grow and a need for an independent baseband digital transport network in parallel with the HFC transport to carry the traffic back to the internet POP.

Another key issue with data going forward is the nature of user demand. It is likely that a small subset of users, consisting of SOHO and telecommuters, will demand and be willing to pay for a massive increase in bandwidth -- into the 10 Mb/s range. The dilemma is finding these customers, servicing them without knowing a payoff and coping with the inevitable churn -- all while not penalizing the cost of servicing the rest of the customer base.

Unlike business users, who can justify a dedicated fiber deployment, these high-end residential customers require a network which can overlay fiber rate bandwidth for deployment as needed.

Finally, for video services, operators must either increase their network transport capacity to move video streams from a central point, or distribute video servers throughout their network. The distributed server model is viable for content which is

common to most users, e.g. "top-ten" movies. Some operators however, already see an opportunity to bring a much wider array of content to the on-demand or personalized model. To enable this business model the cost per stream of optical transport required needs to continue downward in cost.

THE SOLUTION -- NEW CARRIER CLASS ARCHITECTURE

To address these issues, operators require two broad capabilities. First, they require improvements in transport capability to enable carriage of all types of multimedia information reliably and cost effectively between their network locations, whether headend or hub. Included in this transport requirement is the ability to consolidate and transport many reverse path signals, which multiply as fiber goes deeper into the network.

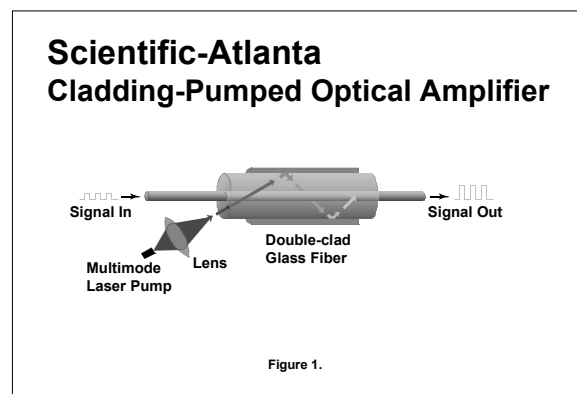
Secondly, the ability is needed to target significant amounts of interactive bandwidth to residential users, including an ability to serve a small percentage of ultra high capacity users cost effectively.

Transport

Operators have traditionally used separate transport networks for voice, data, and video, resulting in inefficiency and a lack of ability to scale independent of shifts in service volumes. Many networks were designed to allow 1550 nm analog optics to be used for video transport. As networks grew, and digital technology came down in cost, many operators shifted toward use of digital video transport, initially using proprietary encoding, and more recently using SONET based systems, such as Scientific-Atlanta's PRISMA DT OC 48 digital transport system. The use of SONET

technology allowed video, voice and data to be carried on a single network. As more and more traffic becomes packet based, however, a channeled data transport system becomes inefficient.

A newly emerging optical transport technique, known as Resilient Packet Ring (RPR), enables an IP-based infrastructure to carry both packet and constant bit rate traffic with the resiliency of SONET based systems and the QOS capability of ATM networks, while eliminating much of the cost and complexity of these technologies. These new IP-over-fiber networks, currently being standardized in the IEEE 802.17 working group, are now entering field trials in MSO networks. They show excellent promise to enable operators to consolidate their existing piecemeal networks over a single network and scale up as traffic increases, independent of service mix. More details of this key network technology can be found at this conference in the paper by Greg Hardy, entitled "IP Transport Technology in the Primary Ring Network."



For those operators using analog video transport, two key enhancements to 1550 nm technology are enabling ongoing improvements in capacity and cost effectiveness of both the headend to hub and hub to node access networks.

The introduction last year of QAM optimized 1550 nm transport has enabled operators to lower the cost of VOD and cable modem transport, and scale sufficient capacity on existing fiber via DWDM counts of up to 24 wavelengths per fiber. Through the use of software selectable pre-distortion operators can enjoy external modulation performance at direct modulation price structures, enabling both a full QAM loading and full distance reach, thereby lowering the cost of VOD stream transport to within the range required for pricing in the typical VOD business case.

A second key enhancement to the 1550 nm optics tool kit has been the recent introduction of a cladding pumped optical amplifier. This technology, as illustrated in Figure 1, utilizes a double clad, all glass ytterbium-erbium-doped fiber to deliver previously unreachable output levels while using highly reliable, extremely low cost technology. Amplifiers using this technology with 27dBm of output power, configured as 8x 17dBm outputs are currently being deployed by a number of operators. The theoretical limit of this technology is 35dBm. By utilizing this technology for the broadcast or common element of the spectrum, operators can cost effectively extend the fiber deeper and deeper into the network.

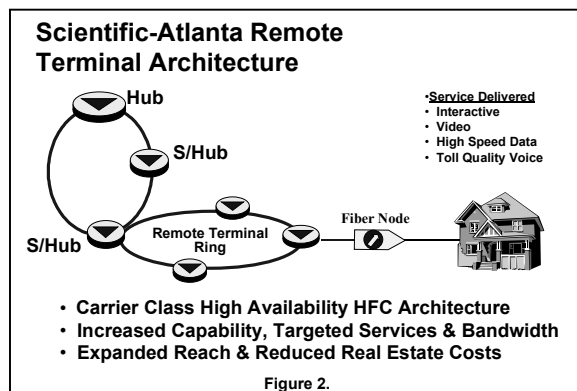
The last key technology being utilized by operators to enhance their networks for high capacity interactive service traffic is baseband digital reverse (bdr). By using Time Division Multiplexing and combining up to four independent reverse path signals and 24 channel DWDM, 96 independent reverse channels can be transmitted on a single fiber. The quality of the transmission can now be engineered by varying the number of bits per sample. This process

eliminates a key weakness in HFC networks: the upstream choke point.

Access

Utilizing the above advances in optical transmission, in both the forward and bdr-based reverse, it is now economically possible to take the fiber all the way to the last active device, producing the so-called “passive HFC” network. Boosting the RF level of the launch amp, made feasible through the elimination of the subtending RF amps, produces additional efficiency.

Given the short coax runs these passive nodes produce, the possibility is raised of running dark fiber in parallel with the coax, for any new construction. By doing this, full service operators can selectively service the



needs of the high bandwidth users, by enabling the fiber, and ultimately migrate more and more users over to the fiber as bandwidth needs grow. More details on the operational aspects of this approach can be found in a companion paper by Don Sorenson entitled “Transition Strategies: Synchronizing Deep Fiber Baseband Access Network Design with Advanced HFC Infrastructure.”

The provisioning of those dark fibers raises a final issue to be considered, however. Even if fibers have been provided

to every termination, an aggregation point needs to be provided deeper in the network than today’s typical hub in order to manage fiber counts.

Transport-Access Connection – The RT Architecture

To address this issue, it is proposed to locate a new network element, the Remote Terminal, deep in the network. This location, connected through a fault tolerant ring to a primary headend or hub, enables groups of fiber deep nodes to be aggregated for distribution of required forward and reverse bandwidth as interactive service traffic grows. This architecture is illustrated in Figure 2.

The RT cabinet is a sited similar to a power supply or splice cabinet. The RT eliminates the need for costly real estate and hub building construction. For maximum efficiency, a hardened optics platform such as the PRISMA II chassis, eliminates the need for costly environmental controls in the cabinet.

Summary

The RT architecture, using all of the new technologies cited in this paper, enables operators to overcome the limitations of first generation interactive HFC technology. By effectively scaling all the way to FTTH, this architecture enables broadband operators to keep pushing the bandwidth limits as they provide increasing value to their customer base.

REFERENCES

LIGHTWAVE; Penwell Publishing; June 2000; Reference Article "Crosstalk in Cable-TV WDM Systems," by Mary Phillips and Dan Ott; Scientific-Atlanta, Inc.

COMMUNICATIONS TECHNOLOGY; Phillips Publishing; February 2001; Reference Article "Deep Fiber Networks: A review of Ready-to-Deploy Architectures," by Donald Sipes and Bob Loveless; Scientific-Atlanta, Inc.

LIGHTWAVE; Penwell Publishing; November 1998; Reference Article "HFC Access Networks Take DWDM in New Directions," by Donald Sipes; Scientific-Atlanta, Inc.

INTERNET TELEPHONY; Technology Marketing Corporation; Using WDM To Cost-Effectively Deploy New Services And Maximize Network Performance," by Jason Shreeram and Donald Sipes; Scientific-Atlanta, Inc.

CABLE TELEVISION; by William Grant; Third Edition; GWG Associates; Amended 1994; Library Of Congress Application number Txu 661-678

"Digitized Return Path: Extending the Applicability of Hybrid-Fiber Coax Architectures for Full-Service Networks," A Scientific-Atlanta White Paper by Farr Farhan, Scientific-Atlanta, Inc.

"The Digital Reverse Path: Cost-Effectively Maximizing Bandwidth For Interactive Services," a White Paper by Scientific-Atlanta, Inc.

CED MAGAZINE; Cahners Business Information; September 2000; Reference Article "Eliminating The Access Network Chokepoint," by Bob Scott and Robert Loveless; Scientific-Atlanta, Inc.

"IP Transport Technology in the Primary Ring Network" by Gregory Hardy; Scientific-Atlanta; Presented at Cable-Tec Expo® 2001; Orlando Florida

"Design, Space and Powering Considerations for DWDM/VOD Deployment in HFC Systems" by Jason Shreeram; Scientific-Atlanta; Presented at Cable-Tec Expo® 2001; Orlando Florida

"Return Path and MTTR Improvements: New Rules, New Tools" by Bob Collmus; Scientific-Atlanta; Presented at Cable-Tec Expo® 2001; Orlando Florida

"Transition Strategies: Synchronizing Deep Fiber Baseband Access Network Design with Advanced HFC Infrastructure" by Donald Sorenson; Scientific-Atlanta; Presented at Cable-Tec Expo® 2001; Orlando Florida

NEW DOCSIS COMPLIANT WIRELESS TRANSIEVER: A LOW-COST, FAST DEPLOYMENT FOR IMMEDIATE REVENUE GENERATION

Ali A. Sayyah
AML Wireless Systems Inc.
260 Saulteaux Crescent
Winnipeg, MB, Canada R3J 3T2

Abstract

This paper describes a low-cost wireless transceiver system that can be introduced as an alternative solution for difficult return-path environments. The design challenges of such a transceiver system are discussed for both CATV and SMATV bands. The lack of a guardband in the CableLabs DOCSIS protocol has been resolved by a novel frequency plan. Also, innovative techniques have been introduced for system automatic gain control and phase locking. Two-way digital testing of the new low-cost transceiver system has shown complete support of the DOCSIS protocol. Finally, testing has demonstrated full Federal Communications Commission (FCC) compliance for the wireless transceiver system.

INTRODUCTION

Emerging from distributing analog video to the perfect pipeline for advanced services like high-speed data and digital video, the cable television industry is faced with ever-increasing demands from its business and residential clients to expand advanced service offerings [1].

Thanks to cable's broadband infrastructure and pervasive networks, the growth of residential high-speed Internet services has been remarkable during the past year [2]. In reaction to such growing demand, CableLabs has introduced two-way system standards by establishing the DOCSIS protocol. However, in some cases the process

of adding return-path to an existing network is hampered by lack of access to support structures, long physical distances or budgetary issues, allowing another ISP to attract the customer by using technologies like DSL or DBS.

This paper describes a cost-effective wireless solution developed to tackle the return-path issue. In addition, the system can be used for extending services to areas where no cable network currently exists. The solution developed by *AML Wireless Systems Inc.* is a full-duplex product incorporating the *ACCESS Transceiver* located in a hub site and an *ACCESS Receiver* placed at the headend facility. The purpose of this paper is to discuss the design challenges of such a transceiver system for both CATV and SMATV bands.

The major design barrier is the lack of a guardband for the DOCSIS protocol between the uplink and downlink frequencies. This issue becomes particularly important for a cost-effective system incorporating a single-polarized antenna at both ends. Two DOCSIS compatible plans for solving these guardband issues have been proposed, meeting the FCC constraints at 13 and 18 GHz frequencies. Another design issue considers upstream phase locking. Because the return-path signal is located in the upper portion of the band, sending the standard pilot-tone would be technically impractical. A lack of phase-lock will result in unwanted phase deviation at the local oscillator, degrading the phase noise and resulting in potential jitters at the receiver. This will be particularly problematic for some cable modems that incorporate CDMA techniques [3]. This solution introduces a

pilot-tone used for phase-locking purposes as well as AGC operation. New block up/down converters have been designed to interface with the cable modem and CMTS equipment at the hub and headend respectively. In order to maximize the available bandwidth, a special filter design has been incorporated to shrink the guardband to less than 13 MHz.

The new cost-effective transceiver system has been tested digitally in full two-way implementations. Downstream results for 256QAM modulation resulted in 35 dB of MER and the upstream link has shown full-band transparency for QPSK, 16QAM and even 64QAM. Additionally, the results of FCC tests (for spurious levels, signal leakage, spectral regrowth and frequency stability) have demonstrated full compliance for the return-path transceiver.

DESIGN CONSTRAINTS

Developing a wireless transceiver, specifically for cable television services, requires careful attention to the legal rules and conditions that apply to licensed band implementations. In particular, a return-path radio connecting to fiber nodes must satisfy the following conditions:

- Transmitting standard VHF channels as specified by ANSI.
- Complying with FCC regulations for fixed wireless services.
- Adapting the constraints imposed by the DOCSIS protocol for two-way data over cable services.

In order to deploy a wireless link, the cable service provider has to obtain a license for operating at CARS-band (13 GHz) or SMATV-band (18 GHz). The following design criteria are applied to the radio:

1. It should transmit the standard channel array without any change or spectrum reversal.

2. The system should comply with FCC regulations.
3. The radio should operate in accordance with DOCSIS requirements. i.e. multichannel video plus up to 256QAM data downstream, as well as QPSK/16QAM data transmission in the upstream path.

Frequency Plan

Based on DOCSIS, the return-path has been specified as 5 to 41 MHz. However 13 and 18 GHz radio links will employ different frequency plans to satisfy the FCC constraints as discussed below:

(A). 13 GHz Return

The FCC limits the bandwidth of the 13 GHz CARS-band to 500 MHz, ranging from 12.7 to 13.2 GHz. A cost-effective solution is made possible by accommodating the upstream channels within the CARS-band. A suggested plan is shown in Figure 1.

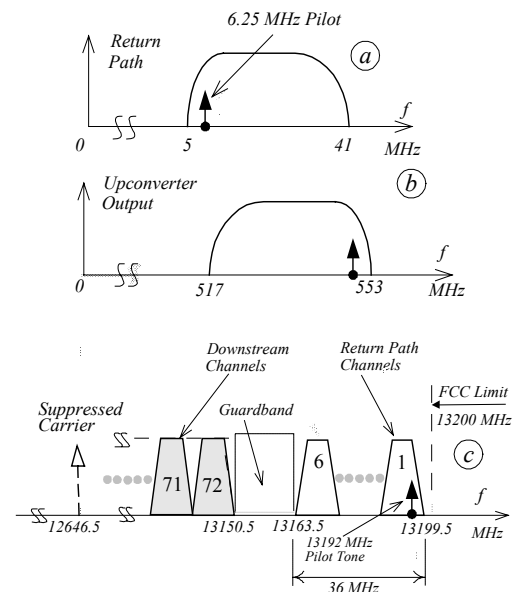


Figure 1. Frequency plan of 13 GHz return-path.

(a) baseband signal with pilot-tone, (b) upconverter output, (c) microwave spectrum.

Figure 1(a) shows the return-path signal. The design feature introduces a 6.25 MHz pilot-tone for AGC and phase-locking purposes. The pilot-tone has been located close to the band edge to avoid interference with the inband carriers. The output of a VHF upconverter module is reversed as shown in Figure 1(b), shifting the pilot-tone to the upper edge of the spectrum. Finally, the upstream microwave spectrum is shown in Figure 1(c), where the return-path signal and its guardband occupy the last eight channels of the CARS-band (Ch.73 to Ch.80).

(B). 18 GHz Return

Similarly, a new frequency plan has been developed for the 18 GHz band, allowing cable operators to deploy a wireless return-path. The design is shown in Figure 2.

The upstream method for 18 GHz is the same as that explained for the 13 GHz design.

Phase Locking

The performance of existing digital modulation schemes is based on phase variation. This makes the modulation sensitive to the phase shift between the local oscillators of the upconverter and downconverter. Generally, the higher the modulation scheme, the more susceptible it is to phase deviation, and consequently the higher bit error rate. The 6.25 MHz pilot is picked up by the downconverter module in order to phase lock its local oscillator with the upconverter, thereby eliminating any phase error. This capability has given the *AML Wireless Systems Inc.* transceiver superior performance in conjunction with highly sensitive CDMA modems.

AGC Operation

The novel frequency plans shown in Figures 1 and 2 cannot include a standard pilot-tone (≈ 74 MHz for CARS and ≈ 73 MHz for SMATV). The pilot-tone is used for AGC operation inside the downconverter module to compensate for received signal variation due to atmospheric changes in the transmission path.

In summary, the pilot tone will add quality to the *ACCESS* system by improving the receiver sensitivity as well as its frequency stability.

LOW-COST RETURN PATH

The maturity of computer technology and the consistent growth of software have made computers affordable to the majority of households. Meanwhile, the advent of Internet protocol has triggered enthusiasm for the computer users to communicate from long distances by exchanging files, data, voice and even video.

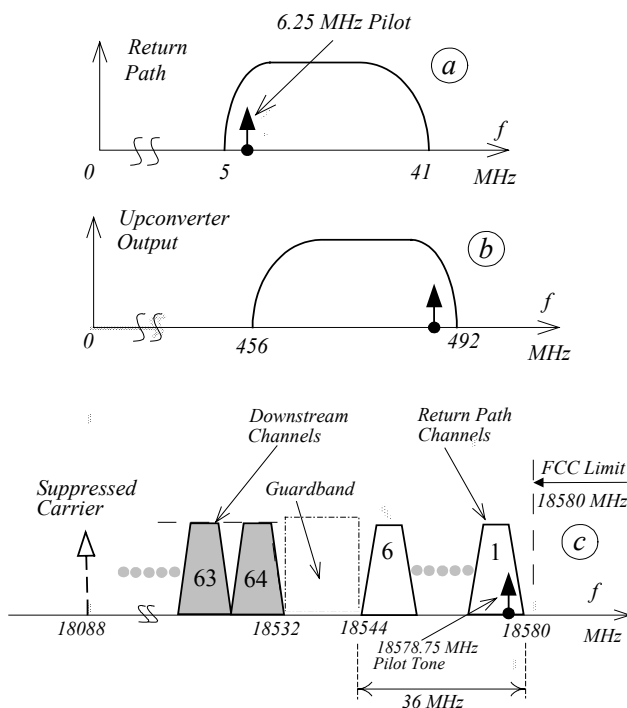


Figure 2. Frequency plan for 18 GHz return-path.

- (a) baseband signal with pilot tone,
- (b) upconverter output, (c) microwave spectrum.

In several cases, the narrow-band, low-speed nature of telephone networks has caused many users to select an alternative Internet service connection. This allows cable service providers an excellent opportunity to emerge into the massive telecommunications market by just adding a return-path to their existing infrastructure, increasing revenue and profits.

The broadband nature of the cable networks, promoting high-speed connections without reliance on the telephone line, has positioned the cable companies as a preferred choice as Internet service providers for both their existing customers and new clients.

ACCESS Transceiver

A low-cost, quickly deployable solution is recommended using new and/or existing equipment. A typical wireless link in a cable network is shown in Figure 3.

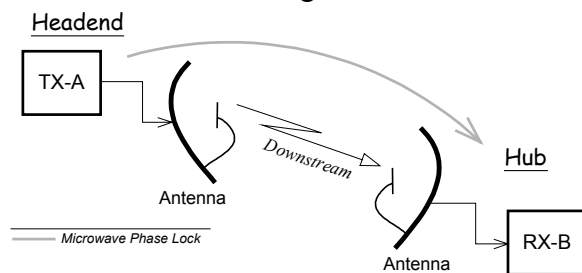


Figure 3. Multichannel radio link in the licensed band

This radio link can be in one of the licensed bands such as 13 or 18 GHz. An upgrade to the return-path is possible by replacing the receiver RX-B with the *ACCESS Transceiver* at the hub site and re-using RX-B as the upstream receiver (URX) at the headend, as shown in Figure 4.

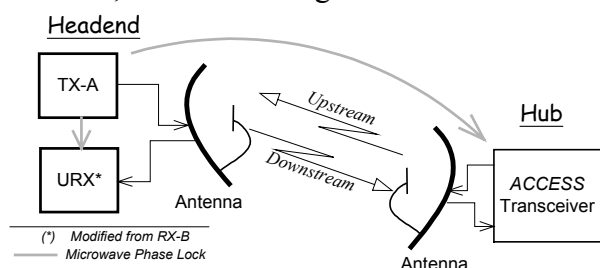


Figure 4. Return-path using ACCESS transceiver

The receiver RX-B will be upgraded in the field by adding the downconverter and performing a few necessary adjustments. A duplex filter assembly to separate the upstream and downstream traffic will be added behind each antenna. Either single or dual-polarized antennae can be used in this upgrade. [4]

Microwave Phase Lock

It should be noted that the 6.25 MHz upstream pilot-tone is used for phase-locking the downconverter unit (at URX) to the upconverter (at the transceiver). The gray lines in Figure 4 show the direction of the microwave phase-lock.

The microwave local oscillator in the transceiver is locked to the headend transmitter (TX-A) using the standard downstream pilot-tone. Similarly, the upstream receiver (URX) must be phase locked to TX-A using a jumper connection, as no standard pilot is sent in the upstream.

A simplified block diagram of the *ACCESS Transceiver* is given in Figure 5, showing the major components.

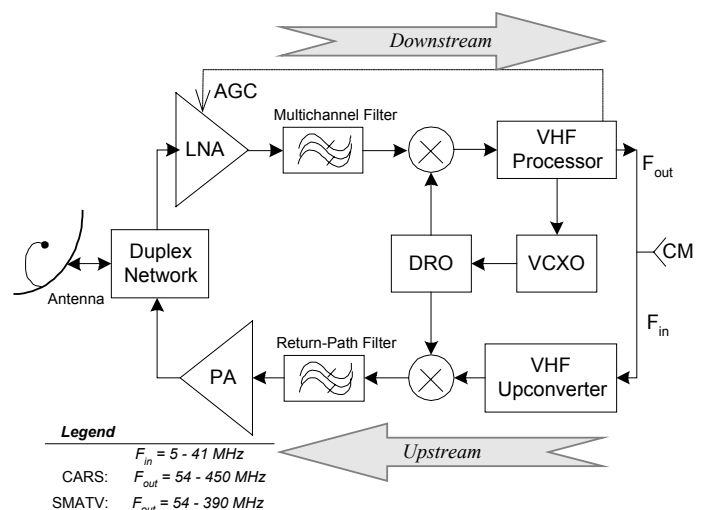


Figure 5. Simplified block diagram of ACCESS transceiver

The input to the transmitter is one or more QPSK/16QAM modulated signals in the range of 5 – 41 MHz. This range will be upconverted to the appropriate VHF frequency as explained. The 6.25 MHz pilot-tone is added to the signal in the upconverter module.

The output of the receiver is a multichannel VHF signal including analog and digital video, as well as data channels with up to 256QAM modulation.

The frequency range will be different for 13 GHz (CARS) and 18 GHz (SMATV) transceivers as explained in Figure 5. All components except the duplex filter network are enclosed in the transceiver housing. The duplex filter network includes bandpass filters and a circulator, inside a weather proof housing with O-ring grooved flanges for pressurization. The duplex filter network is usually mounted behind the antennae.

EXPERIMENTAL RESULTS

The 13 GHz *ACCESS Transceiver* has been tested in order to evaluate the critical parameters required by FCC regulations. The test criteria included spurious level, signal leakage, spectral regrowth, and frequency stability [5].

Bandwidth Overload

A 64QAM test signal occupying 6 MHz bandwidth was applied to the transceiver input. The test signal had a 26 Mb/s bit rate with spectrum as shown in Figure 6.

To verify the effect of the transceiver on signal bandwidth, microwave output has been measured with a spectrum analyzer as illustrated in Figure 7. The transmitter output has been raised to the nominal level (i.e. 6.5 dBm per 6 MHz channel).

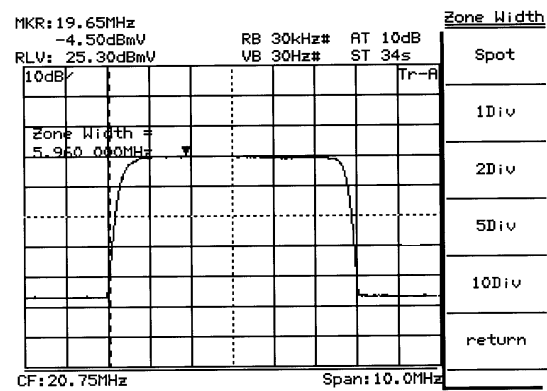
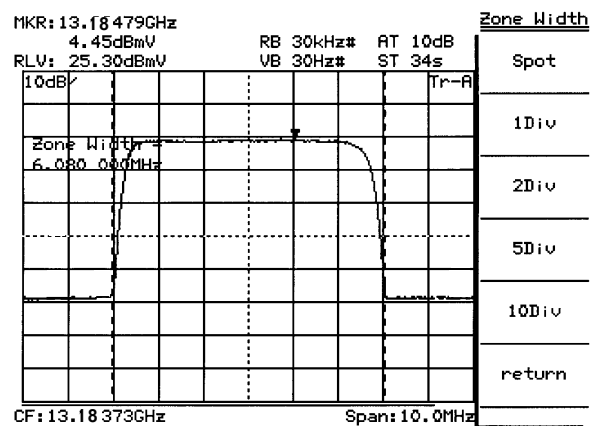


Figure 6. Spectrum of 64QAM test signal as the input



to the transceiver.

Figure 7. Measurement of transceiver output at 6.5 dBm/ channel.

Consistent with FCC regulations, the comparison of Figures 6 and 7 shows that the transceiver does not overload the signal bandwidth.

Spurious Level

To verify the spurious compliance, the VHF LO frequency (558 MHz) has been measured at the transceiver output. This major spurious is only 5 MHz away from the band edge and is difficult to remove.

CONCLUSION

Figure 8 shows the LO spectrum at 13204.5 MHz which is 70 dBc down from the reference level, well exceeding the 50 dBc requirement by the FCC.

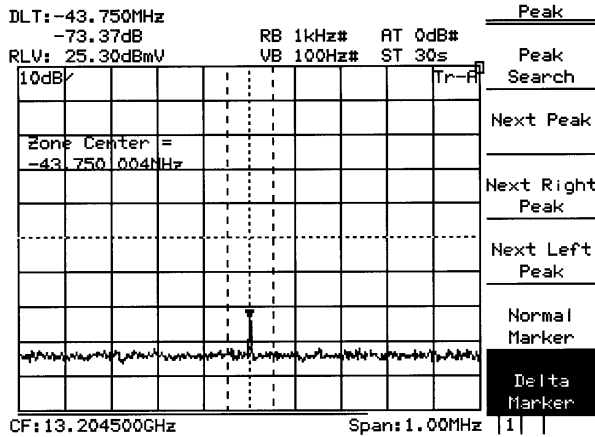


Figure 8. VHF LO frequency appearing in the transceiver output.

Frequency Stability

The performance of this 13 GHz *ACCESS Transceiver* has been monitored throughout a -40 to $+50$ °C temperature range. Figure 9 shows the frequency drift of the microwave carrier, which is well below the FCC limit (5 ppm).

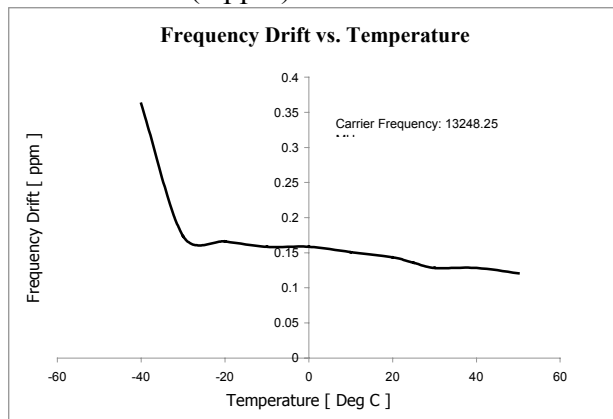


Figure 9. Frequency drift of the carrier frequency vs. temperature.

A low-cost solution has been designed for cable television providers that allows them to quickly deploy a wireless return-path in response to the growing demands for Internet access and immediately increase revenues. Novel frequency plans have been introduced for return-path transport within the 13 and 18 GHz licensed frequency bands. A cost-effective transceiver has been developed to provide full upstream transmission using new equipment or integrating with an existing system. Finally, the test measurements prove the transceiver's performance is in full compliance with FCC regulations.

REFERENCES

- [1] CableLabs, *Cable Television Industry Overview 2000*
- [2] Broadband Intelligence, Inc., *High Speed Internet Competition*, December 2000, p.3.
- [3] CableLabs-issued press release, "*CableLabs Certifies More Modems*", October 20, 2000.
- [4] AML Wireless Systems, Inc., *13 GHz ACCESS Transceiver: Theory of Operation*, February 2000.
- [5] AML Wireless Systems, Inc., *FCC Test Documentation for 13 GHz ACCESS Transceiver*, May 2000.

NEXT GENERATION CMTS CHARACTERISTICS INCLUDING IP MULTICAST

Doug Jones
YAS Broadband Ventures

Abstract

The cable industry is currently faced with upgrading from DOCSIS v1.0 CMTS equipment to DOCSIS v1.1 and PacketCable CMTS equipment. This represents the beginning of a convergence of data and voice onto a single system, what this paper calls a second generation CMTS. This paper discusses an evolutionary path for the convergence of multimedia services onto the Next Generation CMTS.

The Cable Modem Termination System (CMTS) is an integral part of the DOCSIS™ cable data network. The CMTS is essentially the “gateway” between IP services and subscribers. As the network and services change, the CMTS will continue to change too. Operators are encouraged to begin thinking now about the characteristics of a Next Generation CMTS (NG-CMTS) to enable a smooth transition to these products in the future.

The first generation CMTS offered best-effort data service for DOCSIS v1.0. The second (current) generation CMTS is designed for DOCSIS v1.1 and PacketCable™ services. This paper proposes three areas of development needed for the NG-CMTS. These areas include services, form factor, and standardized Application Programming Interfaces (APIs).

With respect to services, the first generation CMTS was designed for data. The current generation CMTS is designed for QoS (Quality of Service) and voice services. A theme of this paper is that the NG-CMTS will take on the role of multimedia processing. Not only are additional PacketCable services, that include multimedia, to be defined, the general trend

in Set Top Boxes is to include a DOCSIS Cable Modem (CM) for interactive services and streaming media. With this in mind, there may be a need for multimedia services in a CMTS. This raises the question about MPEG-2 services moving through a CMTS. These MPEG services may not necessarily go onto a DOCSIS channel, but onto a digital video channel.

IP Multicast is a network capability that enhances multimedia. IP Multicast is the capability to send data from one-to-many recipients, or many-to-many recipients. This approach differs from both unicast, where data is sent from one user to one user, and broadcast, where every user is sent the data whether they want it or not. Both unicast and broadcast data transfer can put increased traffic loads on networks, whereas multicast traffic is selective in only putting traffic where it needs to be. These concepts will be explained in more detail later in the paper.

The form-factor of the NG-CMTS is taking on two distinctive flavors. One camp holds that the NG-CMTS will grow in size and port density, and be highly redundant and survivable. Another camp believes the CMTS will shrink in physical size and will be placed within a fiber node. These opposing views will be explored.

Finally, there are APIs to be supported by the NG-CMTS. APIs can be defined for such purposes as billing, QoS policy, provisioning, etc. If the same APIs are available on all CMTSs, then the operator can create networks using CMTSs from multiple suppliers. If CMTSs implement proprietary APIs, then operators may find themselves locked into a single CMTS supplier.

Vision

With respect to data, the old business was a fast web browsing service for early adopters. The new business will include moving packets, both IP and MPEG, around on networks that will support millions of every day users. The new packet-based services will require more integration and a more defined management plane. The NG-CMTS will be at the heart of this network.

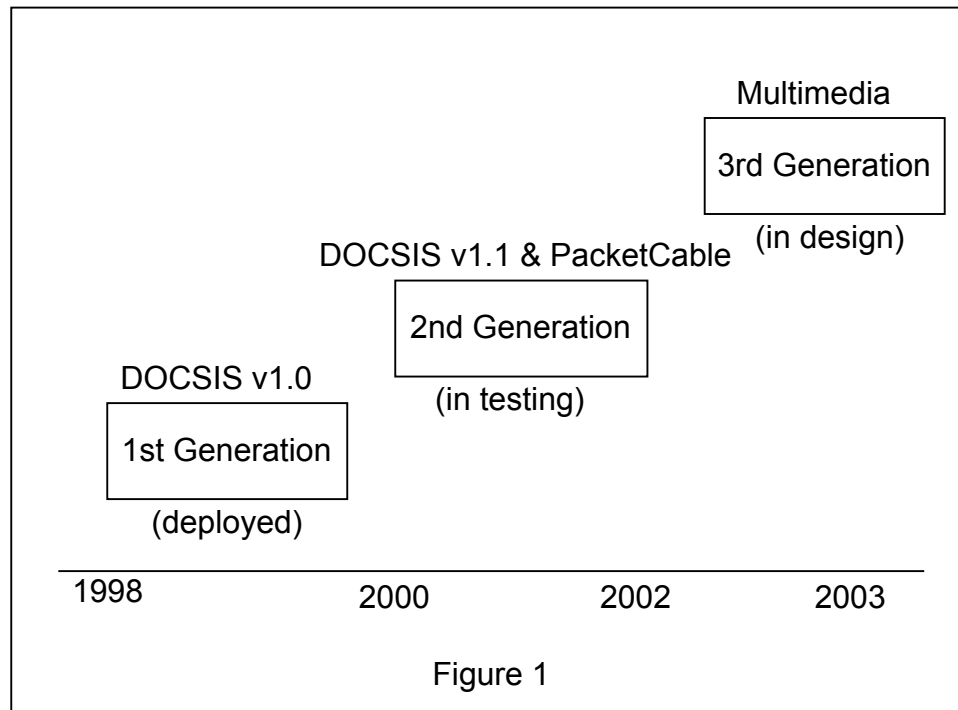
DOCSIS used to be just for high-speed data. Now DOCSIS will be the platform for offering all IP services over cable, including voice and interactive television. In order to meet these needs, the CMTS must continue to evolve. As shown in Figure 1, the

suppliers, DOCSIS would not be where it is right now.

These products were designed to quickly enable a standard high speed data market. The primary service was best-effort data for email and web surfing. These devices have proven to be robust and scalable, supporting interoperability of CMs from many suppliers. These first generation products have proven the concept that mass-deployed cable data service is a reality.

Second Generation CMTS

The second generation CMTS products offer several design advances over the first generation products. One of the largest



industry is well into the second generation CMTS development, in fact, these devices are already being tested at CableLabs in both the DOCSIS and PacketCable programs.

First Generation CMTS

The first generation CMTS products should be commended for the fine job they did in making DOCSIS an accepted world-wide standard. Without the efforts of these

drivers of this has been the success of the first generation products; there are 4 times as many suppliers in the second generation market and this drives competition and innovation. In addition, PacketCable has defined services that require higher reliability and stringent Quality of Service (QoS), both of which have contributed to new capabilities on the CMTS. Finally, the DOCSIS v1.1 specification has defined new features and

functions for the second generation CMTS, including dynamic services, account management, IP Multicast, SNMPv3, certificate-based authentication, etc. In addition to the features already mentioned, these second generation CMTSs will differ in their power consumption, port density, foot print, cabling options, etc.

Due to these new guidelines, the second generation CMTS is more sophisticated and feature rich than the first generation CMTS. With the large number of suppliers in the market, operators have many new choices to consider. With the second generation CMTS, operators should be realizing just how central the CMTS is to their services plans.

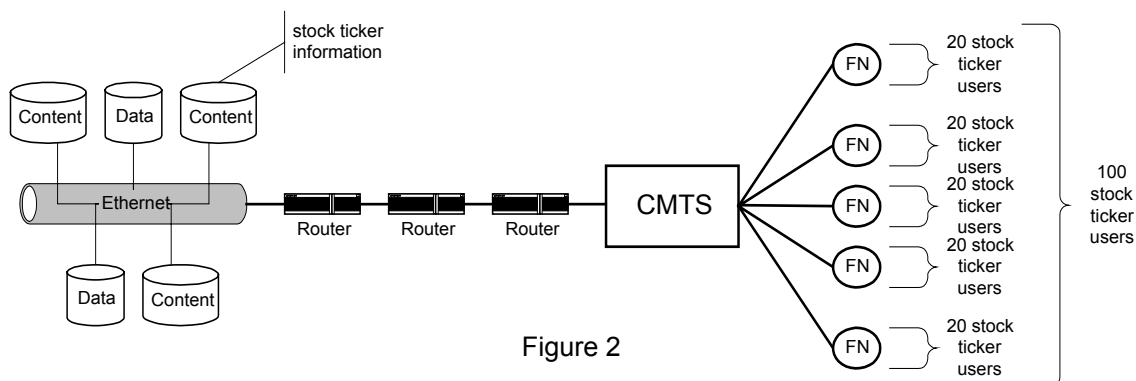


Figure 2

Second Generation CMTS - IP Multicast

The second generation CMTS, as part of the DOCSIS v1.1 specification, must implement defined rules for IP Multicast support. Specifically, DOCSIS v1.1 places rules on the CMTS for implementing the Internet Group Membership Protocol (IGMP). IGMP is the underlying protocol that allows IP Multicast services to work.

In DOCSIS v1.0, several operators experimented with IP Multicast services. However, without defined IGMP support in the CMTS, these initial attempts at IP Multicast service were technically successful, but not scalable. Having defined rules in the CMTS will allow operators to now push

forward with real market offerings using IP Multicast.

Potential services include streaming headlines, stock tickers, and digital audio. While these are low bit rate services, if many subscribers access them with unicast service flows, then the amount of bandwidth consumed grows with each new user. However, if the services are multicast, then the IGMP support in a second generation CMTS ensures that the content only flows along network segments where necessary. Multicast users on a segment each “tune” into the single flow that is present on that segment.

With a unicast model, if one user on every fiber node subscribes to the service, then there is one instance of that data on each node. If 20 subscribers on each node want that service, then there are 20 instances of that data on each node. The difference with multicast implemented is that if 20 subscribers on each node want the service, there will only be one instance of that data on each node, and each user can access it. In a sense, that data is shared by all the users. This is a powerful tool that will conserve bandwidth not only on the backend network, but on both the forward and return paths as well. In order for IP Multicast to be most efficient, each router in the network, from the content source to the CMTS, must be IGMP enabled.

Figure 2 can be used to illustrate the benefit of IP Multicast. In Figure 2 with unicast IP traffic, the server farm would need to source 100 individual streams to feed the 100 users. This traffic would burden every LAN segment, switch, router, and CMTS between the data and the end user. Since there is only a single source for the stock ticker, the individual unicast streams would all contain the same information; hence, there would be a lot of redundant information on the channel. This consumes bandwidth.

With IP Multicast, and depending on router configurations and headend combining, the server farm would need to source only 1 data stream to feed the individual fiber nodes. The CMTS would then replicate that stream to each downstream needed to feed the fiber nodes. The savings in bandwidth is readily apparent.

Next Generation CMTS

Given that the second generation CMTS has had many advances, operators may question if there is a need for a NG-CMTS. Now we get to the interesting discussions.

With the NG-CMTS, operators will truly begin moving into converged services. That is, all IP devices connecting to the cable plant will get their services through a CMTS. Right now, IP is thought of as a service that allows subscribers to do email and web surfing. With the NG-CMTS, operators should also consider moving completely to IP as the method to control devices that connect to cable plant. Clearly IP is the world standard for internetworking and with the adoption of both DOCSIS and PacketCable, the cable industry is gaining more and more experience with IP services. Since these two services run over IP, the cable operators will continue to be developing their IP expertise. Expanding the CMTS to include the management and control of all IP devices on the cable network seems to be an evolutionary step.

The benefit would be a reduction of both operational and capital cost for the cable operator. Currently for a headend that offers multiple services and connects to multiple backend networks, there will be separate racks of equipment for each service and network. With services converging over a NG-CMTS, there is the potential for having “fewer boxes” in the headend. In addition, rather than training technical staff to support 3 separate sets of equipment, each a complex technology in and of itself, operations support can begin consolidating around a single technology, the NG-CMTS. In addition, the innovation provided by a multi-supplier network will benefit operators both from the standpoint of feature availability and cost.

Next Generation CMTS – Features

Given that the second generation CMTS is designed for both data and voice, it is possible to envision the NG-CMTS taking on the role of multimedia and video processing. This would make the NG-CMTS responsible for both IP and MPEG services. Adding MPEG services to the CMTS could include such features as:

- Cherry picking
- Rate remultiplexing
- Conditional access

The MPEG processing in a NG-CMTS could relate to either streaming media or broadcast quality VOD content, or both. The choices will be made by the operators. But, having MPEG processing in the CMTS follows from a theme in the proceeding section. If all devices connect to the CMTS for both IP services and control, this would include a next generation set top box (or home gateway) that not only offers IP and MPEG services to the user, but also gets its signaling and control information over IP. CMs and Media Terminal Adapters (MTAs) already get signaling and control information over an IP interface to the CMTS, so the industry is heading in that direction. These types of standard devices, based on the

DOCSIS protocol for carrying IP over cable networks, are being developed now and should be available in 2002.

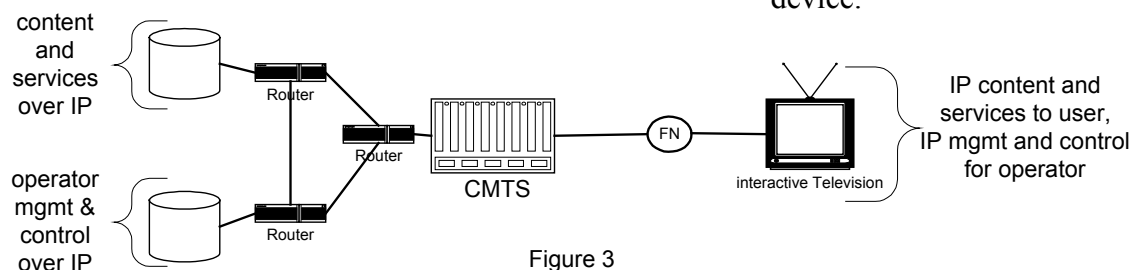


Figure 3

Figure 3 shows the concept of the operator and service providers embracing IP for control and services, respectively.

In addition to considering MPEG services, the NG-CMTS will also need to consider dynamic traffic management services, including tools for both controlling overall available bandwidth and how that bandwidth is apportioned to various services and users by QoS mechanisms. DOCSIS v1.1, in the second generation CMTS, introduced the basics of these services, and based on learnings, they will probably be fine-tuned in the NG-CMTS. These tools will be needed not only for subscriber-facing services, but also looking into the back-end network to support both other Internet Service Providers (ISPs) and Alternative Service Providers (ASPs) that are connecting to or providing services over the cable network. These ISPs and ASPs will have contractual agreements with the cable operators for certain levels of service, defined by Service Level Agreements (SLAs). The SLAs will need to be verified and enforced. This will very likely be an important business consideration for cable operators.

Next Generation CMTS – APIs

For first generation equipment, DOCSIS v1.0 placed very few requirements on the CMTS. This was an operator decision to encourage innovation among the CMTS

suppliers. The following were two main areas of specification:

- Standard IETF Management Information Bases (MIBs) for the management of the CMTS as an IP device.

- RF requirements to maintain the integrity of the cable network.

Second generation CMTS equipment, under both DOCSIS v1.1 and PacketCable, have had additional requirements placed on them that include:

- Dynamic Quality of Service
- Network time synchronization
- IP Multicast Support
- Certificate-based authentication
- Additional standard MIBs
- SNMPv3
- Etc.

As can be inferred, the CMTS is becoming a more complex piece of equipment.

Continuing to define standard APIs for the CMTS is an important consideration. These APIs should be defined to allow operators to implement standardized control and management functions across the CMTSs. With more and more suppliers entering the CMTS market, operators can either take the route of specifying additional standard APIs, or be prepared to have suppliers solve these needs in a proprietary fashion. While proprietary solutions are sometimes how features are first developed, with operator input these features can be

migrated into the specifications. With the number of suppliers in this space, and the competition and innovation going on, these APIs will appear on CMTSs. But being locked into a single supplier in a highly competitive IP services market may be detrimental. By contributing these interfaces to the specifications, suppliers will make it more feasible for operators to deploy CMTS equipment across diverse networks.

For the NG-CMTS, new APIs that allow a move toward more common management, security, and provisioning should be considered. The issue of concern, again, is having standard APIs on CMTS equipment to allow operators to procure equipment from multiple suppliers and retain their operations and management frameworks.

Operators should still have the capability to define their own APIs in order to best meet their business needs, but as these interfaces mature, or as issues are identified, migrating the API to the specifications should be considered.

Next Generation CMTS – Location

This interesting question revolves around where the NG-CMTS will physically reside. CMTS configurations are available to fit a variety of options to allow operators to design networks as they choose. Possibilities include placing the CMTS in either an environmentally controlled facility such as a distribution hub, or moving the CMTS closer to the subscriber by placing it in a fiber node.

Clearly operators are placing the CMTS in a controlled facility today. However, several suppliers have demonstrated equipment that will fit the CMTS into a fiber node. As these new products mature, operators will have another choice in designing their cable data networks. The “CMTS on a pole” concept is still very new, and operators are encouraged to monitor these developments over the next year.

Locating the CMTS in a fiber node allows the operator to place processing closer to the subscribers. This partitions the DOCSIS network into smaller segments, similar to how making node sizes smaller partitions the RF network. This has the benefit of a more distributed architecture, making an individual CMTS a smaller point of failure. On the other hand, the CMTS would be placed in an environmentally harsh environment, and may be more accessible to tampering.

Regardless, a CMTS in a fiber node is an interesting development and the industry should pay attention to how it develops.

Summary

The evolution of the CMTS does not stop with DOCSIS v1.1. Operators should debate the functions of the NG-CMTS as its quite possible all data, voice, digital video, and device control services may some day be moving through this piece of equipment.

The CMTS will continue to evolve in terms of the services it supports, the APIs that are needed to control the network, and the physical form factor and location of the CMTS. In order to prevent having to do hardware upgrades of CMTSs at a later date, operators should consider the NG-CMTS services they want now and figure these into their services plan for the future.

References:

1. Miller, C. Kenneth, Multicast Networking and Applications, Addison Wesley, 1999.

Author Contact:

Doug Jones, Chief Architect
YAS Broadband Ventures
400 Centennial Parkway
Louisville, CO 80021
voice: (303) 661-3823
fax: (303) 661-3830
doug@yas.com

NON-INVASIVE NETWORK MANAGEMENT

Bruce F. Bahlmann
Alopa Networks

Abstract

Broadband operators have extremely limited visibility into the health of Hybrid-Fiber Coax (HFC). Their limited visibility is due to several factors including the complexity of their transport medium, type of network, and sophistication of their back office. To navigate these factors, a more non-invasive approach to managing HFC is needed. This paper will present such an approach that is completely scalable, offers sophisticated location logic that enable one to quickly locate the most common types of HFC outages, utilizes a mere fraction of the bandwidth required by other network management efforts, and will fully integrate into existing top-level network management systems without additional stand alone status monitors.

Introduction

Broadband operators increasingly find themselves in a tough spot. Which is trying to figure out how to support and manage new services that are nearing deployment or have already been deployed. HFC being a complex animal as it is but now having to manage multiple services running on-top HFC leaves most broadband operators scratching their heads.

The fast paced deployment of new services is usually managed with a brand new Network Operations Center (NOC). These NOCs are quickly staffed with as many people (most internal) as it takes to provide 24-hour coverage 7 days a week. The result of promoting installers and plant operations personnel provides the broadband operator with increasing challenges in terms of organizing people and equipment to perform the minimalist amount of

monitoring of the HFC and its increasing number of services – training is a huge issue.

The scarcities of technical employees skilled in network operations to design, build, and run broadband operator NOCs contribute to increasingly ugly statistics on building successful network management organizations. Over 70% of the attempts to initiate network management resort in failure.

Considering the cards stacked against broadband operators, they generally succeed in providing a minimal level of network management over HFC. However, the areas they face the most challenge are providing in-depth multi-service network management support, expanding their visibility beyond their backbone and hubs – down to their End of Lines (EOL), and developing more advanced associations (correlations) between related events. While most broadband operators would say they would like to explore these areas there is this problem regarding their obvious lack of commercially available tools as well as experienced network operations people to use them.

Background

Managing HFC is far from an exact science. In fact, HFC is so crammed with various types of technology, transmission media, and content that it is extremely tough to keep operational. It is also difficult for broadband operators to have visibility all the way down to their EOL – the point(s) at which each HFC node terminate. As a result of much of the HFC being invisible and complex, Broadband operators end up picking and choosing which components of HFC represent the most critical and monitor them with whatever tools are available. Basically the rule of thumb is the more customers that share

the same HFC component and/or transport the more critical it becomes (thus their current focus of the backbone and hubs).

Unfortunately, there are still not many tools available for managing HFC (the actual fiber nodes) – certainly not many that don't represent a completely new and often stand alone system. Since the last thing broadband operators are looking for is yet another monitoring system to drop into their already over crowded NOC, not many companies have gained traction with broadband operators for their HFC management products. This has left much of the market of managing HFC open. In fact, broadband operators to this day still do not have a cost effective way to manage all the way down to their EOL. The key here is cost effective – as some vendors have provided solutions for EOL monitoring but they are extremely cost prohibitive and not all together realistic to deploy operationally.

The focus of this article is to explore ways that broadband operators can gain more visibility into their HFC nodes including an example EOL monitoring system.

HFC Monitoring Challenges

The reliability of any network management system is directly dependent on the extent that it reaches out to all of its network elements. Customer Premise Equipment (CPE), Cable Modems (CM), Set Top Boxes (STB), and Media Terminal Adapters (MTA) represent attractive customer located network elements for NOCs to reach out and verify HFC health and availability. These installed customer network elements are placed throughout the network, represent no additional cost, and provide increasingly useful status and health information to a traditional Network Management System (NMS). However, traditional NMS rely on actively polling network elements to collect operational status of the environment where these network elements reside. Thus, current efforts to use customer

network elements along with a traditional NMS fall short of the mark because they rely on active polling of these network elements. The method of active polling suffers from several issues that will be explained.

Scalability is one of the most obvious issues. Essentially, the sheer numbers of customer network elements can reach a point where it impacts the frequency that a single application can poll them on a regular basis within a timeframe that is worthwhile. As a result, the frequency that network elements would be polled by a traditional NMS would be increased so as to allow all these elements to be polled without impacting the network performance that is trying to be measured and monitored.

Because there are so many network elements to poll this also prevents one from being able to obtain much of any detail from each network element. The more information obtained from each element the greater the time it takes to collect this information, the greater each request impacts the usable bandwidth of the network it seeks to monitor, and the greater the impact on the network element performance it is attempting to use to monitor it.

Not being able to poll frequently presents another problem if using a traditional NMS. This is because the reason one polls network elements is to determine their current status and look for potentially service degrading behavior. This does not work well with HFC as it changes constantly. While much of these changes are tolerable (hardly noticeable operationally) the changes that can indicate much more serious problems are brewing also happen sporadically. Since a traditional NMS can only poll periodically it is likely to surmise that it will not be able to capture (or detect) these sudden changes and thus be unreliable in determining much more than trivial (on/off) status of the HFC.

Another problem with active polling is that unless similar network elements are polled

together (or within a reasonable time frame) the information gathered is useless across all similar network elements. For example, all network elements on a network can span several HFC nodes (i.e. they are combined). Unless all network elements are polled by-HFC-node and within a reasonable timeframe the information gathered may only indicated that something is potentially wrong with one of the HFC nodes. However, since none of this information can be collated by-node, the resulting data is unreliable and only marginally useful. NOCs that actually poll network elements on various nodes usually resort to managing subsets of network elements, if at all, so they can only look at very small-controlled samples of the customers.

Traditional NMS applications are best suited to manage network elements with static Internet Protocol (IP) addresses and are not capable of managing network elements with dynamic IP addresses. NOCs that do manage customer network elements must re-map the IP addresses of these network elements with every renumbering of the network. Since NOC resources are at a premium, this often results in fewer instances of monitoring customer network elements on each HFC node. This practice only leads to increasingly less monitoring of the HFC.

Monitoring customer network elements using a traditional NMS pays a heavy price on ones network because it uses Simple Network Management Protocol (SNMP) get. An SNMP get requires a NMS to send a question to a network element somewhere on the network. Each question in SNMP terminology corresponds to a specific Management Information Base (MIB) located in the network element's operating environment. Essentially, each network element maintains a wealth of MIBs each of which corresponds to some configuration or operational data stored locally on the network element. Depending on what is asked (i.e. which MIB(s) is/are requested in the get) the network element responds to the request by determining all the answers to these

questions and then sends back a reply. As a result of this transaction, the network between the NMS and the network element pays twice for this transaction – once for the get and again for the reply.

Since not many customer network elements (e.g. a CM) actually have a need to communicate with the outside world they often fall off most routing tables. What this means is when some other system decides that what information they have is all of the sudden useful it must blaze a trail to each element from a networking perspective. This process of blazing a trail involves creating routing table entries for each network element. Routers that provide connectivity for each network element outside its network must re-learn about these network elements before communications can flow between the NMS and the network element. Each routing table entry requires the network element's router to Address Resolution Protocol (ARP) its physical address. ARP allows one to determine the mapping of IP address to Media Access Control (MAC) address (also called physical address). Once this table entry is created in the router it is able to relay packets to the network element. Albeit, this process is extremely quick (even in network time) this latency across all subscribers only further contributes to the inefficiency of using traditional NMS on network elements.

There is also this issue of determining the correlation between the network element and its associated real world information. One can usually derive certain things given pieces of information. For example, knowing ones IP address and subnet mask other information about the network can be derived. Likewise, knowing ones phone number or first and last name one can determine where one lives. However it is impossible to derive relationships between dissimilar or unassociated things. While there exists ways to complete these relationships they currently go beyond the capability of traditional NMS. For example, typical billing and customer care systems can

associate network elements with real world customer information. However, NMS do not provide hooks into such systems. Instead, these associations must be built manually – a very tedious and unmanageable process.

The NMS is also the wrong tool to manage HFC because it must monitor customer controlled network elements. That's right, all these network elements belong to the customer (or at least an increasing number of them do with the advent of retail CMs, etc.). So many of its requests to these network elements will not go through – such is the unpredictable nature of a customer-controlled device. As a result, it must actually ignore many of its responses because they will come up empty (or non-responding). In the early stages of deployment of CMs (or any other new technology) only a handful (if any) of network elements may exist on each node or Cable Modem Termination System (CMTS). What this means is that most (if not all) network elements may be down and this could actually represent 'normal' operating conditions.

Focus Areas

There are two general areas one needs to focus on when monitoring HFC. These general areas are:

- Monitoring HFC health
- Monitoring HFC EOL

There are "solutions" that claim to address both of these areas using one technology, but to do this right these areas actually represent two drastically different approaches that no one system can sufficiently achieve.

For example one system uploads tables of information specifically from their CMTS to provide some visibility into the HFC. However solutions like these are highly proprietary and will not work in a multi-vendor environment. They also provide only limited health monitoring, can impact the performance of their

CMTS, and do not address EOL monitoring. Keep that in mind when looking at network management systems that claim to address both of these areas.

Since there are many reasons (several were discussed previously) why traditional NMSs are not up to the task of performing reliable HFC monitoring one must explore uncharted territory to achieve the visibility needed to manage/maintain high service quality.

Uncharted Territory

The most promising technology that can facilitate scalable visibility into HFC is actually quite old and extremely well established – the use of SNMP traps. SNMP traps provide a non-invasive way of monitoring the health of ones HFC plant because they are non-solicited. Network elements capable of SNMP traps are first configured to look for certain events/conditions (e.g. some threshold is reached), and then inform their configured trap host when these traps conditions are met. SNMP traps provide the following attractive features:

Completely scalable – Since SNMP traps use individual network elements to capture its information there is no need for any one application to perform direct polling on network elements. Once more, SNMP traps can be easily directed to any number of applications further distributing the load of handling all the traps across multiple servers. The distribution of this load can be easily added over time and more importantly in conjunction with the number of network elements.

Minimized network element performance impact – SNMP traps actually minimize the performance impact as compared with that of direct polling. That is because the network element only communicates an SNMP trap if a threshold is met rather than continually respond to SNMP requests. In the mean time the network element merely examines these thresholds along side its normal function.

Network elements monitor more states – Unlike direct polling which must carefully optimize what it request from network elements, SNMP traps can look at wider variety of MIBs. MIBs whose data only seldom changes are not good candidates for direct polling. However these make terrific SNMP traps as they send extremely valuable information about the network element at the time it occurs. This optimizes the collection of SNMP trap responses and allows applications receiving these traps to be overall more responsive to the needs of these network elements.

Network elements report independently – SNMP traps allow network elements to report their information independently rather than as the result of being polled directly. Reporting independently allows network elements to report the moment their thresholds have been met. This also allows applications receiving these traps to correlate multiple responses so as to determine the extent and severity of these traps.

Tolerant of dynamic IP addresses – Since SNMP traps come from network elements they are inherently tolerant to changes in the IP address on the network element. While this does force the application receiving the traps to be cognizant of each network elements' current IP address, this problem is much more manageable than the challenges that would exist to provide this functionality in a traditional NMS.

Half the bandwidth cost of SNMP gets – SNMP traps (like SNMP gets) communicate over something called the User Datagram Protocol (UDP) in networking world. However the SNMP trap does not require an acknowledgment from its destination. Therefore once a network element sends a SNMP trap it is done – no waiting around for some type of reply. From a networking perspective the use of UDP is an extremely efficient means of communicating. Note however that UDP requires a fairly reliable network to operate properly because if the network drops packets (i.e. it is unreliable) it

may very well drop the SNMP trap message. In this case it is a self-fulfilling prophecy in that SNMP traps are used to make the network that much more reliable.

Tolerant of delays in routing – SNMP traps don't care about how long they take to reach their destination. Likewise, their originating network elements also don't care how long they take to reach their destination. As discussed previously, once they are fired out of the network element the communication is over as far as the network element is concerned – which goes on about its business (no matter what is left for the SNMP trap to negotiate to reach its destination).

Address correlation between element & customer – SNMP traps do not in themselves provide any real correlation between the network element and the customer. However what they do provide is a means of translating these traps at the application responsible for receiving the traps. At this point network information and real world information can easily meet. This particular area is where focus is needed to build the necessary relationships between the network element and the customer it represents. When these two pieces of information meet the result is an extremely powerful database capable of advanced reporting, troubleshooting, and modeling.

Tolerant of being customer controlled – SNMP traps can easily withstand having a customer connecting and disconnecting as well as powering up and down the network element. In the event the network element can communicate over the broadband medium its information is transmitted along side others. Should it be shut down or disconnected, it will not participate in the collective monitoring of the network health. In this way it makes no difference what so ever if the network element is on or off. Having it on would be great but if it is shut down it doesn't break anything or cause any false alarms.

SNMP traps form the basis for providing an excellent HFC health monitoring system. Combine this with an intelligent trap collection application that can be distributed and you have a fairly cheap means of monitoring HFC health.

Monitoring HFC Health

Monitoring HFC health is an evolutionary process. Just remember that managing HFC is far from an exact science. To do this right one needs many different sources to accurately track its health. One source can come from CMs, one from MTAs, yet another from STBs, and so on. It is important that when one allocates spectrum to these services that it strategically selects frequencies across the entire spectrum. Sticking to only those frequencies that are known to be good defeats the purpose of using customer network elements to monitor the health of ones HFC. If various customer network elements are not positioned across the available spectrum the use of SNMP traps will not be able to provide sufficient sampling to determine overall HFC health and should therefore not be used. Essentially this would merely provide visibility to a narrow portion of the overall spectrum. One actually needs several reference points (at least 3) across the entire spectrum to provide any kind of reliable HFC health monitoring.

Several ancillary benefits may be achieved out of a comprehensive use of SNMP traps to monitor the health of ones HFC. Some of these include:

Less reactionary network operations – Given the NOC now has visibility to significant changes on the HFC it can direct resources to make repairs before these changes become overly noticeable to customers.

Significant individuals can be more closely observed. Very important people (VIP) as well as past trouble makers can be more carefully watched without drawing attention as would be the case if the were modeled within an NMS.

Exploring more extensive use of SNMP traps will elevate their importance in the eyes of standards bodies and CableLabs. This will result in more specific trap requirements in future network elements geared more specifically towards monitoring network health.

Monitoring HFC EOL

Unfortunately, monitoring HFC EOL represents a totally different process than monitoring the health of HFC and is relatively void of cost effective solutions. This is because EOL monitoring concerns itself with connectivity and availability where as monitoring HFC health concerns itself mainly with reliability. Albeit these are somewhat related, connectivity and availability are distinctly different from reliability. Essentially monitoring HFC EOL insures the entire physical plant is available (no breaks, outages, etc.) all the way down to its EOL. Consequently the types of messages sent by EOL network elements are quite different than those sent by HFC health type network elements. This actually represents a relatively new area for broadband operators in terms of deploying/using extensive EOL monitoring. In fact, there are not many commercially available EOL monitors that will provide a cost effective solution to this problem.

Today's EOL monitors are extremely expensive (around \$200 US), proprietary, and require some type of line voltage where they connect to the HFC. These requirements force many broadband operators away from the technology (even though they really would like to have it). Instead broadband operators seek to use other means of observing the HFC but none of which provide them with the same kind of visibility. The requirement of an available line voltage is also highly restrictive as this is not always available near EOLs – certainly not every EOL.

One cost effective EOL means of monitoring may be through the use of a miniaturized CM. The concept here is to drive chip technology to reduce the footprint of a CM down to something

that would fit in an enclosure no bigger than a line filter. In this case, a slimmed down CM is placed in a line filter. It is then powered by the 90v square wave that flows down the HFC for telephony. As a result one can have an operational standalone CM capable of sending and receiving information on the network with one slight exception. Instead of responding as a normal CM, these EOL CMs will provide some additional functionality. This includes the ability to perform predictable chatter. Chattering is a process where by the EOL network element periodically talks to its host. This predictable chatter allows the system to determine possible outages.

Any network element that is past due (not heard from) would need to be followed up using some basic form of direct polling. In this case direct polling is used to speed the resolution of the event. In other words, is the network element down or is some link down. Once this information is obtained one can forward this information to the broadband operator's NOC for analysis and resource assignment.

However new the concept of monitoring EOL is to the broadband operator its benefits outweigh the barriers needed to make this component a necessary part of the overall network management system. Yet broadband operators' general lack of interest in this component is likely the result its current lofty price tag. Until these components can get within the \$20-30 US range (or lower) there use will not enter the main stream.

Bruce Bahlmann
Alopa Networks
bahlmann@bigfoot.com

OPENCABLE APPLICATION PLATFORM ARCHITECTURE

Allen R. Schmitt-Gordon, Ph.D.
Cable Television Laboratories, Inc.

ABSTRACT

The OpenCable Application Platform (OCAP) is a software middleware layer that resides functionally on top of the Operating System of a OpenCable terminal or receiver. It provides an interface enabling application portability. A fundamental requirement is that applications written for OpenCable be capable of running on any network and on any hardware platform, without recompilation.

The OCAP specification is built on the DVB MHP 1.0 specification with modifications for the North American Cable environment that includes a full time return channel. A major modification to the MHP is the addition of a Presentation Engine (PE), that supports HTML, XML, ECMAScript. A bridge between the PE and the Java Execution Engine (EE), enables PE applications to obtain privileges and directly manipulate privileged operations.

INTRODUCTION

The OpenCable Application Platform (OCAP) is part of a concerted effort, called OpenCable(TM) by the North American cable operators (MSOs) to provide the next generation digital consumer device, encourage supplier competition and create a retail hardware platform. A cable receiver that can be provided at retail must provide interoperability and portability of content and applications across networks and platforms, and be geared towards the full range of interactive services. Current devices are network specific and operate proprietary software that is not portable across platforms or networks. With OCAP, the applications are written primarily to a middleware

(software layer between the operating system and the application software) API so that a single application can be deployed across the full range of OpenCable host devices available at retail. Such applications include:

- Electronic Program Guide (EPG)
- Impulse Pay Per View (IPPV)
- Video On Demand (VOD)
- Interactive sports, game shows
- E-mail, Chat, Instant messaging
- Games
- Web Browser: Shopping, Home banking
- Personal Video Recorder (PVR)

Another essential requirement is that the middleware be secure and robust. Stability in the Cable terminal or receiver is imperative and resets are not acceptable. The middleware must also simplify content development, through a publicly available application programming interface. Finally, the OCAP must be operating system and hardware agnostic. That is, OCAP will specify only one profile. Applications and content will be written to only one profile.

Background

This development of the OpenCable middleware was initiated in September, 1999 through the RFP process, 16 vendor companies submitted proposals by October 15, 1999 for a middleware architecture that would enable such portability. Review of the submitted proposals was completed by the end of December of 1999 and several middleware vendors were selected to refine the architecture and develop the OCAP specification. Initial specification development was started in January of 2000 and development teams comprised of the

selected vendors, members of the Cablelabs staff and visiting engineers were assembled.

The RFP was worded such that the specific components of the middleware architecture were not specified but left open to the respondents. The architecture that was chosen by the technical team was a middleware that was comprised of two parts: a Presentation Engine (PE) and an Execution Engine (EE). The PE generally was composed of an HTML engine and ECMAScript. The EE included a virtual machine. This architecture is shown in Figure 1. It shows that native applications are supported as well as application written to the middleware via the OCAP interface.

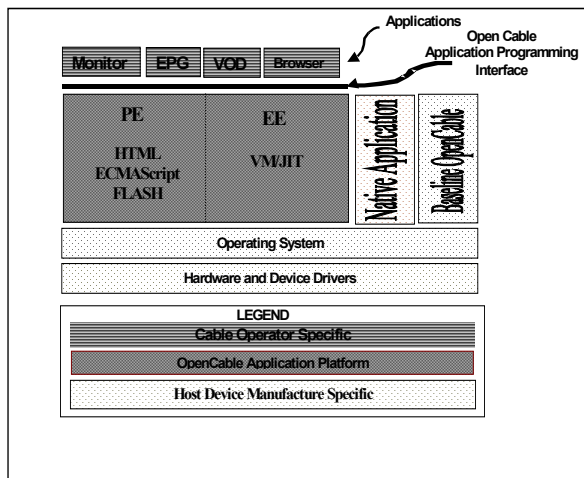


Figure 1

One of the key elements to the development of the OCAP architecture has been the licensing agreement with Sun Microsystems. Sun has provided pertinent portions of the Java API specification and related IP to Cable Labs. Anyone implementing the OCAP specification may implement the Java API without obligation to Sun. Cable Labs will incorporate Sun Technology Compatibility Kit (TCK) as part of Open Cable compliance test suite, and only Sun Java licensees may use Java branding. This agreement enabled the

OCAP team to specify the Java Virtual Machine (JVM) and JavaTV as the fundamental components of the EE.

In order to expedite development of the OCAP specification, it was necessary to utilize existing standards and architectures, as much as possible.

ARCHITECTURAL COMPONENTS

Rationale for Middleware

The current software architecture in general usage, requires that all applications be written to the specific Operating System running on the hardware platform being used. In an environment where cable receivers are leased and are derived from a limited number of suppliers, this model works to some extent. It does not enable platform or network portability. Applications such as EPG, VOD, mail, etc are compiled to the application programming interface determined by the operating system and associated hardware. This is illustrated in figure 2.

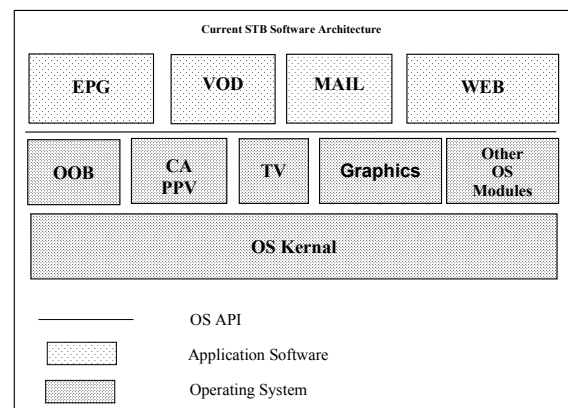


Figure 2

By imposing a middleware layer, that abstracts the functionality of the OS, hardware devices and network interfaces, such as shown in figure 1, a variety of applications can be written that will run on any platform. Such applications will run, without recompilation, as would be the case for applications written

in C or C++. Careful design will enable network portability as well.

Presentation Engine

The PE enables the use of tools that have been widely used for internet content. The PE renders declarative content such as graphics, text, animations and audio based on formatting rules in the PE itself and formatting instructions such as the markup language.

Its primary components consist of HTML 4.01, XHTML 1.0, CSS 1 and CSS 2, and ECMAScript 3. In addition, advanced TV is supported through Macromedia Flash and Plug-ins to access other Web content formats.

Unless the PE has many extensions added to it, certain functionality will be lacking. Most notable would be tuning. In order to facilitate this functionality and not circumvent the inherent security of Java as well as provide a common mechanism for managing receiver resources, the middleware architecture will include a bridge between the PE and the EE.

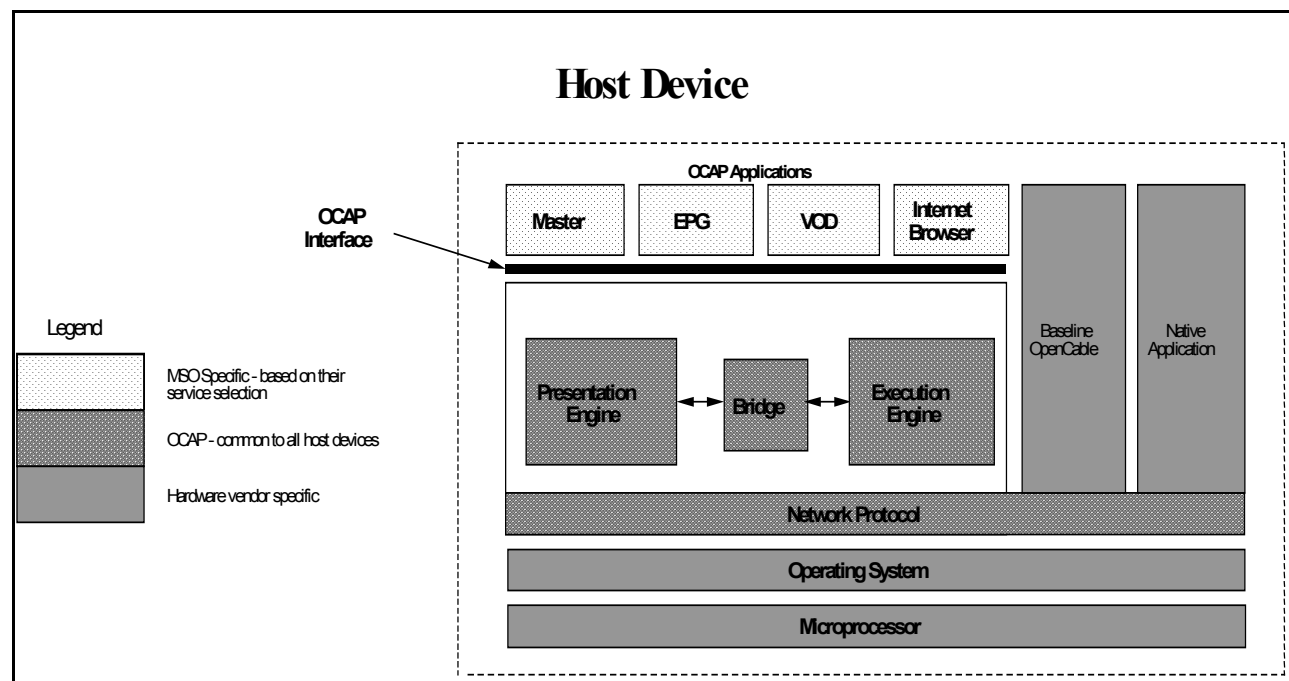


Figure 3

For example, as part of the EE that is accessible via the bridge, JavaTV offers a common point of control and management of various system resources that includes tuning. Thus, a PE application will access device resources through the bridge. This ensures that device resource contention is managed through a common control point for a PE and

EE application. that are vying for the same resources.

Execution Engine

The EE provides a general application programming environment for networking, file I/O, graphics, etc. The OCAP EE (with MHP) provides a full TV application environment.

Java is more portable than C or C++ and provides a platform independent byte code. The EE provides a platform independent set of APIs.

As a starting point in the development of the EE, the DVB-MHP 1.0 specification was chosen to build upon in order to expedite the specification development process. OCAP specific extensions have been added to cover those elements such as a full time return channel, application management, and resource, service information.

The major elements of the EE are control of application management through the pJava APIs, control of application management through the pJava APIs, service information and selection through the JavaTV API's, media control through the JMF, broadcast data through the MHP DSMCC APIs. In addition, the EE provides network management and IP data access and extensions from HAVI and DAVIC and DASE.

A fundamental feature of the EE utilizing Java is that security is built into the architecture from the ground up.

Bridge

In order to enable the browser to take full advantage of the resources in the receiver, the design of figure 1 was expanded to include a bridge between the EE and the PE. This is also shown in figure 2. The bridge permits access by ECMAScript applications to the Java Class Libraries and Java programs access to the DOM files. Thus a full programming environment is available for application written for the EE or the PE.

Through the use of the bridge to extend PE functionality beyond what would normally be possible for a browser, the use of plug-ins is minimized. Plug-ins would normally be written in native code, so that many ports of each plug-in would have to be maintained.

Conclusions

The OCAP architecture offers the widest range of support for high quality, attractive, TV centric applications and display of content currently being proposed or offered. It offers a very high degree of portability and uniformity for content display as well as offering a platform for the broadest possible range of application support. The OCAP architecture and certification process ensures security and robustness.

References

Kar, M.L., Vang, S., and Brown, R., Architecture of Retail Set-Top Box Application Platform for Digital Cable Network, Proc. ICCE, June, 2001.

Zundel, J-P, Emergence of Middleware in Home Telecommunication Equipment, IEEE Communications, June, 2001.

DVB MHP 1.0 -- TM220 3 Rev.4 CM195 Rev.1 SB28(00)07, DVB Multimedia Home Platform Revision 16

ATSC DASE AEE, Doc. T3-530 09, Feb 2001, Rev 1.

Havi Level 2 User Interface, section 2.5.2, <http://www.havi.org>

Documents relating to the PE including DOM, CSS, HTML can be found at <http://www.w3c.org>

Documents relating to the Davic specification can be found at <http://www.davic.org>

RETURN PATH AND MTTR IMPROVEMENTS: NEW TOOLS, NEW RULES

Bob Collmus
Scientific-Atlanta, Inc.

Abstract

While broadband networks have seen many advancements and technical innovations in recent months, one in particular stands out as a rule changer – that of digital return.

The return or upstream path of HFC network architectures has seen a dramatic change over the last 18 months as a new technology – known as baseband digital reverse has been deployed in increasing numbers. While the key technical benefits of a digital return band are increased return bandwidth, extended optical reach, and the elimination of performance barriers common to analog transmission, the use of baseband digital return further opens up opportunities for improving performance while lowering mean time to repair (MTTR).

A primary benefit here is realized by reducing overall equipment requirements. Architectures that are reduced in component count are made possible with this technology and this paper explores and demonstrates by example how transmitter counts have been lowered by 50%, how fiber counts have been reduced by 50% and how optical amplifiers have been eliminated altogether.

As our industry trends toward more and more interactive services, the demands on the reverse path increase dramatically. Not only are network operators concerned about whether there is sufficient bandwidth to address all the requirements, but they are equally concerned about the performance and reliability of the return path. Baseband Digital Reverse (bdr™) has emerged as a

key enabler for delivering all three essential return path components – bandwidth, performance, and reliability. This paper will explore how bdr delivers on this promise by looking at the technology itself, the applications it enables, and the architectures it supports.

The Technology

Since Scientific-Atlanta first introduced bdr at the SCTE in 1999, much has changed. We are currently on our third generation of the products and are continuously finding new applications for the technology. Because much has been said over the past two years, I will not spend a great deal of time reviewing the technology in detail, but do feel it is important to establish a foundation for the benefits it can provide.

At its simplest, Figure 1. below shows a single bdr link. The 5 to 42MHz RF return is converted from analog to digital by a high speed (100Mbps), 12 bit converter. The resulting 12 data streams are sent to a serializer to produce a 1.25Gbps data

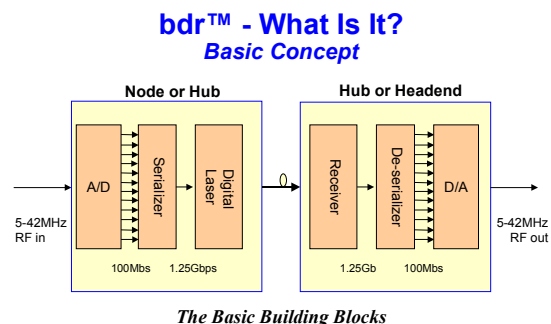


Figure 1

stream. This signal is then transmitted upstream using a low-cost, digital laser. At the receive end, low-cost optics and

electronics are employed to convert the signal back to its analog format for processing.

From a performance and reliability standpoint, bdr technology exhibits several key attributes:

1. excellent carrier-to-noise performance - 12-bit analog to digital converters provide “DFB-like” performance; the noise power ratio (NPR) curves (Figure 2) perhaps best demonstrate the differences in FP, DFB, and bdr technologies.
3. distance insensitivity – digital link performance does not change in performance over distance; this simplifies installation and eliminates performance changes when systems operate over back-up paths which are typically longer.
4. optical reach – digital receiver sensitivity is greater than its analog counterpart; since no degrading of performance is experienced over distance, we are easily able to deploy links over 80km without the need for optical amplifiers.

NPR Performance

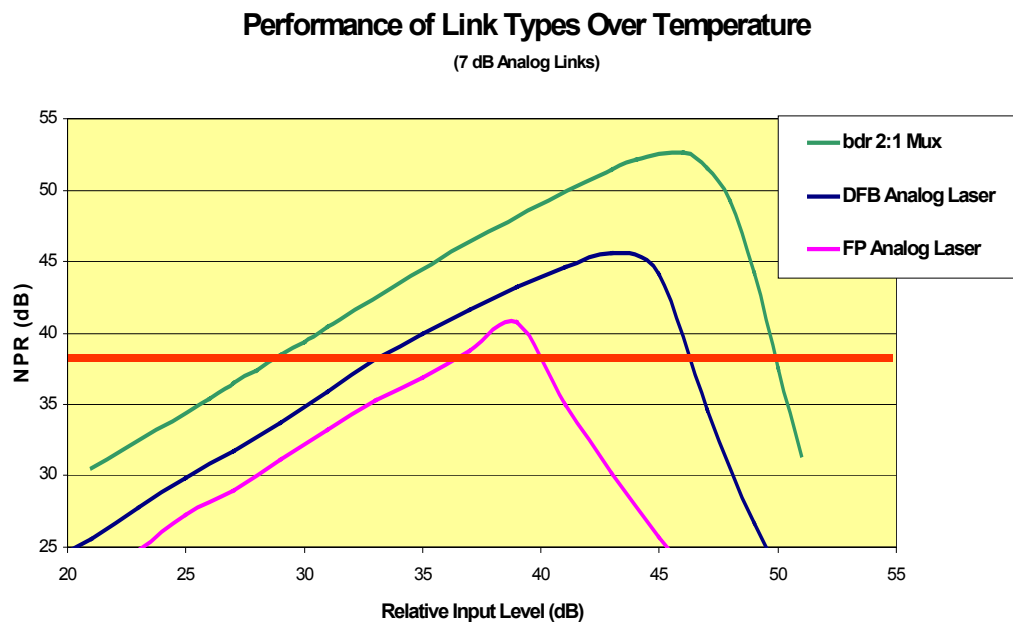


Figure 2

2. temperature stability – digital lasers are minimally impacted by temperature; therefore, negligible changes in system performance are seen over wide temperature ranges, a significant benefit over analog lasers.

The Applications

At the recent SCTE Emerging Technologies conference, it was clear that, while 2001 was

the year for the emergence of VOD, 2002 promises to be the year for Voice over IP (VOIP) to be deployed in large scale. Couple VOIP with other emerging applications that are symmetrical in bandwidth usage, such as Napster (when it surfaces again), video teleconferencing, and streaming video, and the tremendous growth expected in the small office, home office market, it is clear that the reverse path of HFC networks is about to undergo a significant increase in usage.

selecting distributed feedback (DFB) lasers or bdr to achieve the required performance needed to support these higher-order modulation schemes. The temperature instability and “noisy” condition when unmodulated limit the application of the analog FP lasers.

Beyond supporting QAM transmission to increase bandwidth out of cable modems, operators are looking to increase bandwidth

Time Division Multiplexing Basic Concept

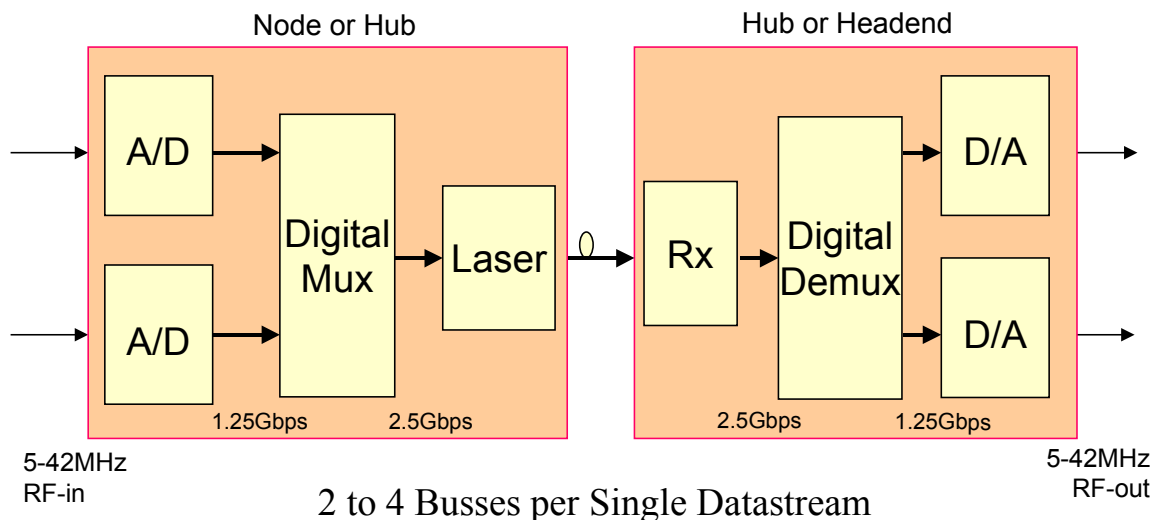


Figure 3

To meet the increased bandwidth challenge, many operators are minimally planning to adopt the DOCSIS 1.1 standard for their IP services. Equipment designed to meet this standard make use of 16QAM and 64 QAM for reverse transmissions. However, these efficient schemes require significantly better (6 and 12dB, respectively) carrier to noise performance than QPSK - the current standard approach. As a result, many operators are moving away from low-cost analog Fabry-Perot (FP) lasers and are

throughout their networks. Installing more fiber is an obvious solution, but a costly one. Technologies that “mine” more bandwidth from the existing plant are needed. This is the important role of bdr and dense wave division multiplexing (DWDM). Unlike analog, bdr enables the use of time division multiplexing (TDM) to more cost-effectively transmit multiple channels over a single fiber. An example of a simple 2:1 TDM application is shown in Figure 3. In short, two 1.25Gbps data-streams are summed together and transmitted at 2.5Gbps using higher speed optics.

Scientific-Atlanta has also announced a more complex 4:1 solution that utilizes digital signal processing to compress four channels of data into a single 2.5Gbps data-stream.

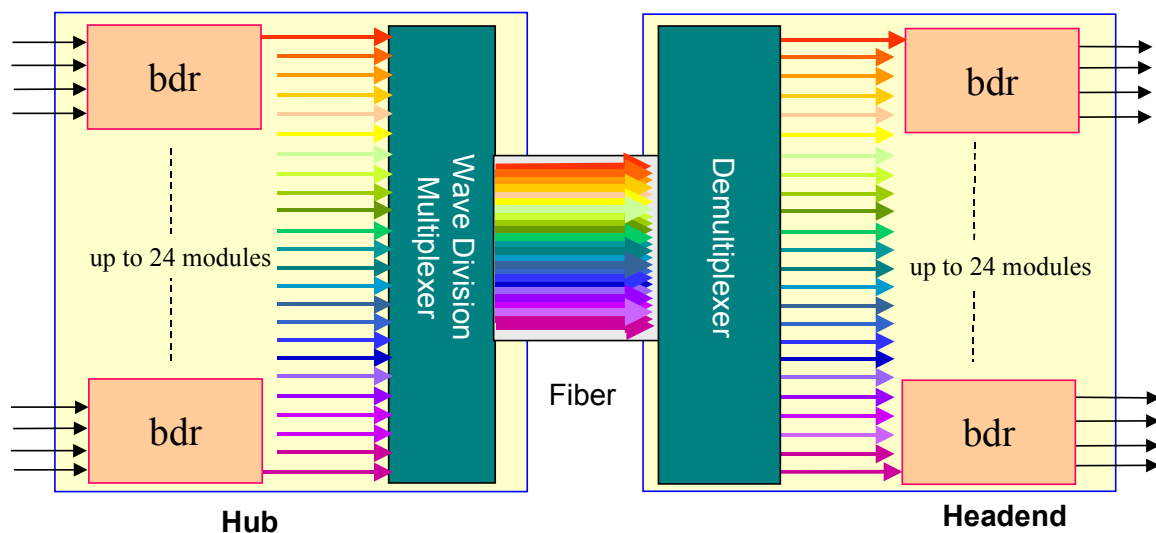
DWDM and bdr technologies working in tandem provide even greater fiber efficiency. DWDM enables multiple optical signals (Scientific-Atlanta has shipped 24 channel systems to date) to be transported over the same fiber.

The Architectures

Certainly, bdr can effectively support the variety of architectures that traditionally have been deployed by network operators. However, bdr and DWDM have opened up several new options that have proven to be more cost-effective to build and, potentially, more reliable to maintain. The focus of these new architectures is a move to a more centralized processing approach.

The centralized processing architecture (Figure 5.) allows the operator to locate

Time Division Multiplexing with Dense Wave Division Multiplexing



Up to 96 Busses on a Single Fiber

Figure 4

The TDM dimension of bdr multiplies the number of DWDM channels by up to four, enabling 96 return paths over a single fiber, as shown in Figure 4.

processing equipment – CMTSs for data and voice IP-based services, QAM and QPSK modulators and QPSK demodulators for digital video services, HDTs for circuit-switched voice services, and servers for status monitoring and control services – in primary hub or headend facilities. Under more traditional architectures, this equipment would be located in the

Centralized Processing Architecture

Key Benefits: Real Estate Savings & Resource Allocation

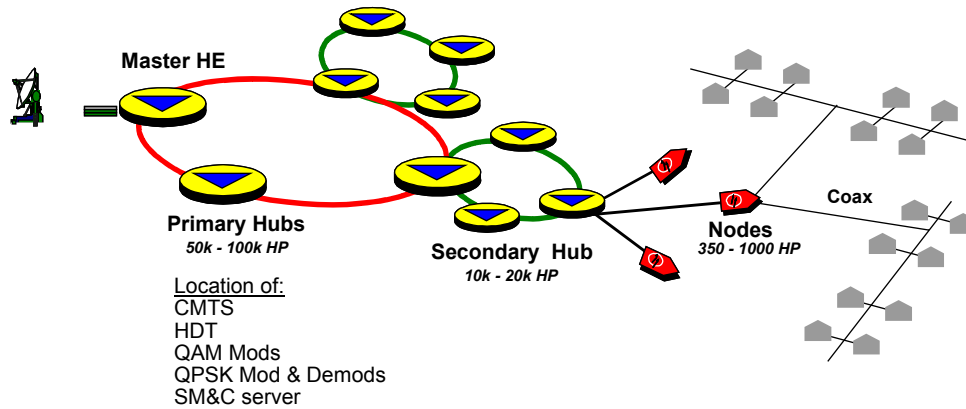


Figure 5

secondary hubs scattered throughout the serving area. bdr is essential to this architecture because it provides long optical reach and minimizes fiber usage. The obvious key benefit to this approach is the cost savings associated to real estate and personnel. However, there are other

Centralized Processing Architecture

The centralized architecture can be accomplished with an analog approach, as well. However, the cost and reliability of this solution are not attractive. Shown in Figures 6 and 7 are the architectural block diagrams for the analog and digital reverse solutions. Approximate costs for the return

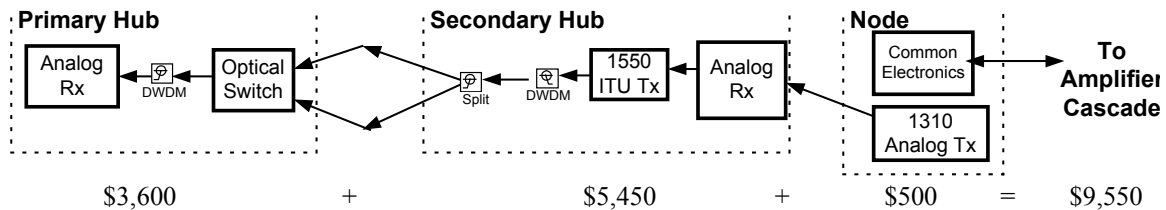


Figure 6 - Analog Solution Block Diagram

important reliability enhancing benefits, such as improved sparing, more complete utilization of equipment, more redundancy options, and quicker response times.

path specific components are shown for each solution. The network shown does not leverage the TDM aspect of bdr. When this is taken into account, even more significant cost savings result.

The reliability block diagrams are shown below in Figures 8 and 9. Included in the diagram are the calculated outage minutes per subscriber associated with the various network components.

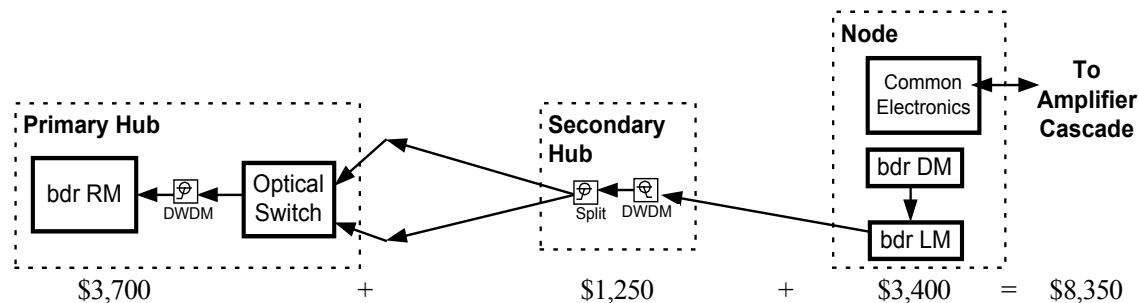


Figure 7 - bdr Solution Block Diagram

Figure 8 - Analog Solution Reliability Diagram

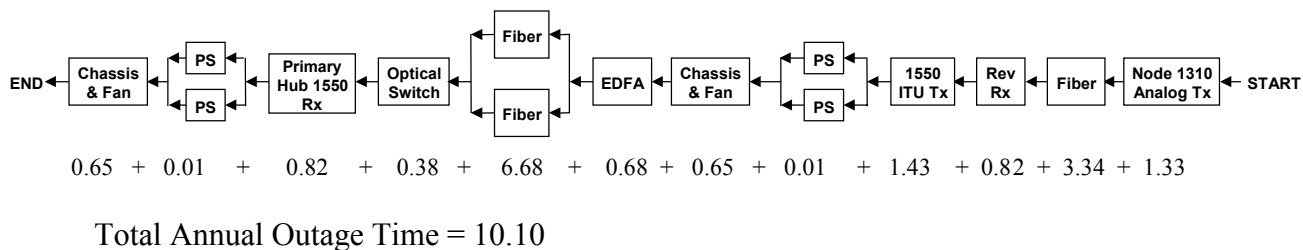
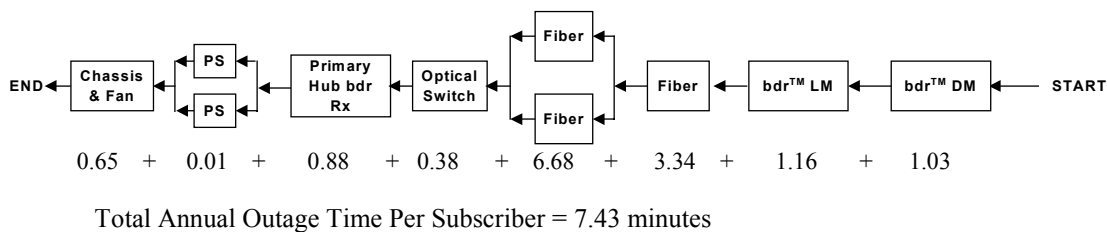


Figure 9 - bdr Solution Reliability Diagram



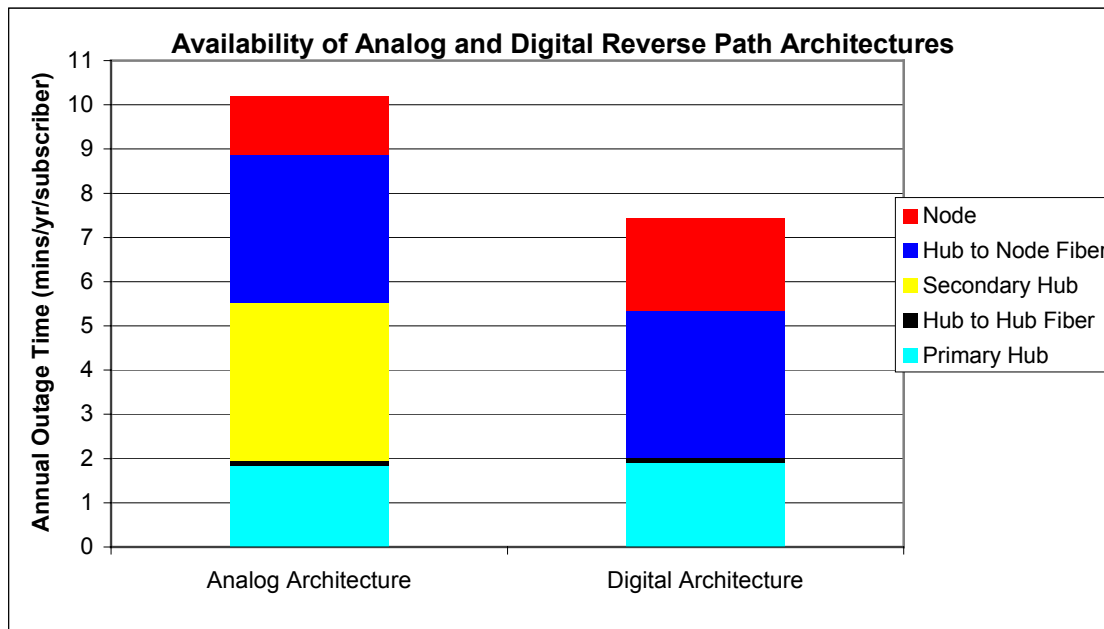
These outage time estimates are calculated using Markov Models, Bellcore Reliability Predictions and field failure data in a four step process developed by Scientific-Atlanta. The results of this modeling are shown in the table of Figure 10.

Figure 10 – Modeling Results

Location	Description	SA Part#	Type Redun	PCP - Predicted MTBF (hrs, M1C1 QL1)	Enhanced MTBF (hrs.)	MTTR (hrs.)	MTTC (hrs.)	Availability	Average Annual Outage Time (Min.)
Analog Reverse Path Primary Hub to Secondary Hub to Node									
Primary Hub	Dual Reverse Rx, P2 - Video SA Conn	716480	2	150,234	1,606,896	24	2.50	0.99999844	0.8183
Primary Hub	Optical Switch, P2 - SA Conn	714470	2	326,456	3,491,759	24	2.50	0.99999928	0.3766
Primary Hub	P2 Chassis & Fan	594300	2	189,449	2,026,338	24	2.50	0.99999877	0.6489
Primary Hub	Prisma II Chassis PS	589254	2		833,333,333	24	2.50	1.00000000	0.0016
Fiber	40 KM fiber	NA	1/4+		314,941	4.00	NA	0.99998730	6.6800
Primary Route Availability =								0.99998730	6.6800
Availability with Redundancy =								1.00000000	0.0001
Secondary Hub	50 Tx, 6475-6, 750 MHz 110 Ch, Dual Output	NA	2	85,848	918,227	24	2.50	0.99999728	1.4321
Secondary Hub	EDFA, 6476-16T	573070	2	180,323	1,928,727	24	2.50	0.99999870	0.6818
Secondary Hub	P2 Chassis & Fan	594300	2	189,449	2,026,338	24	2.50	0.99999877	0.6489
Secondary Hub	Prisma II Chassis PS	589254	2		833,333,333	24	2.50	1.00000000	0.0016
Secondary Hub	Dual Reverse Rx, P2 - Video SA Conn	716480	2	150,234	1,606,896	24	2.50	0.99999844	0.8183
Fiber	20 KM fiber	NA	1		629,882	4.00	NA	0.99999365	3.3400
Node	ASSY, MFLEX FP XMTR SCA	717904	2	92,563	990,050	24	2.50	0.99999747	1.3282
Analog Reverse Path Total								0.99998080	10.0965
Digital Reverse Path Primary Hub to Secondary Hub to Node									
Primary Hub	ASSY,MOD.14BIT PRISMA RCVR MODULE	716157	2	139,225	1,489,145	24	2.50	0.99999832	0.8830
Primary Hub	Optical Switch, P2 - SA Conn	714470	2	326,456	3,491,759	24	2.50	0.99999928	0.3766
Primary Hub	P2 Chassis & Fan	594300	2	189,449	2,026,338	24	2.50	0.99999877	0.6489
Primary Hub	Prisma II Chassis PS	589254	2		833,333,333	24	2.50	1.00000000	0.0016
Fiber	40 KM fiber	NA	1/4+		314,941	4.00	NA	0.99998730	6.6800
Primary Route Availability =								0.99998730	6.6800
Availability with Redundancy =								1.00000000	0.0001
Secondary Hub	NA								
Fiber	20 KM fiber	NA	1		629,882	4.00	NA	0.99999365	3.3400
Node	MOD ASSY, 6940 BDR 2:1 DIGITAL MODULE	712892	2	119,464	1,277,782	24	2.50	0.99999804	1.0291
Node	ASSY,1560.61NM 6940 BDR XMTR 2.5GBPS	713312	2	106,408	1,138,135	24	2.50	0.99999780	1.1554
Digital Reverse Path Total								0.99998586	7.4347

A comparison of the bdr and analog approaches can best be depicted in a graph format as shown below. Notice the large contribution of the Secondary Hub analog components to the overall outage time. This link in the chain is totally eliminated in the bdr solution, greatly improving overall system reliability and simplifying any maintenance requirements.

Figure 21 – Availability



Remote Terminal Architecture

Scientific-Atlanta recently introduced a variation of the centralized processing architecture discussed above. The Remote Terminal (RT) network (Figures 12 & 13) utilizes a third ring to reliably push fiber closer to the customer, thereby significantly increasing the available bandwidth per home passed.

unprotected fiber links. The network and reliability block diagram (Figure 14) depicts the longer link, yet shows that overall outage time per subscriber is actually reduced over the centralized architecture approach examined previously.

Figure 12 - Remote Terminal Architecture

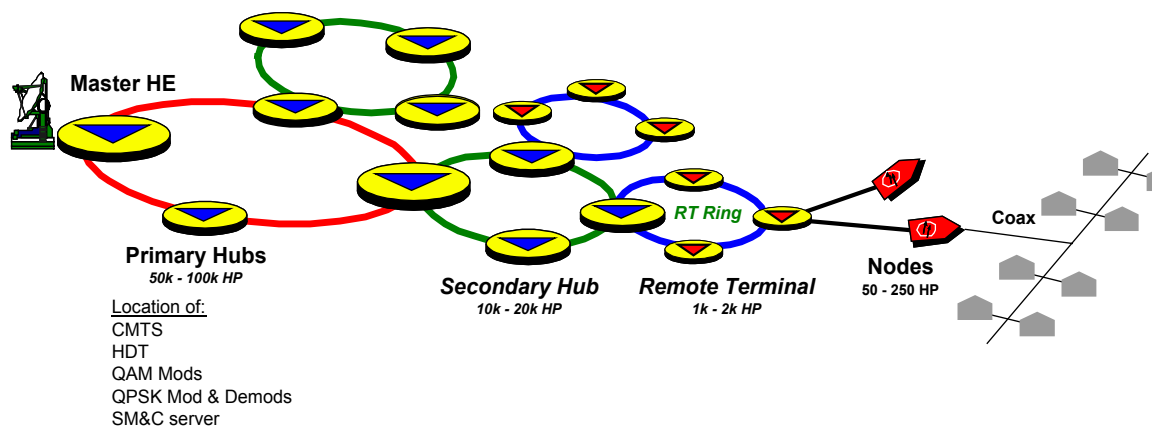
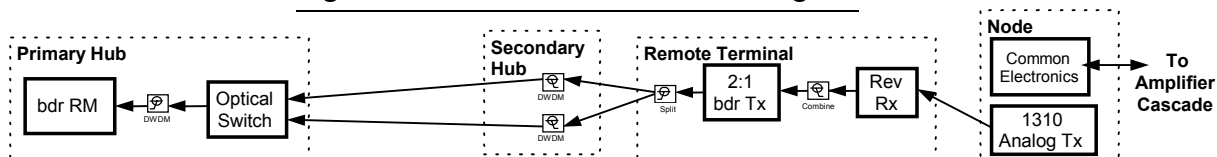


Figure 13 - RT Architecture Block Diagram



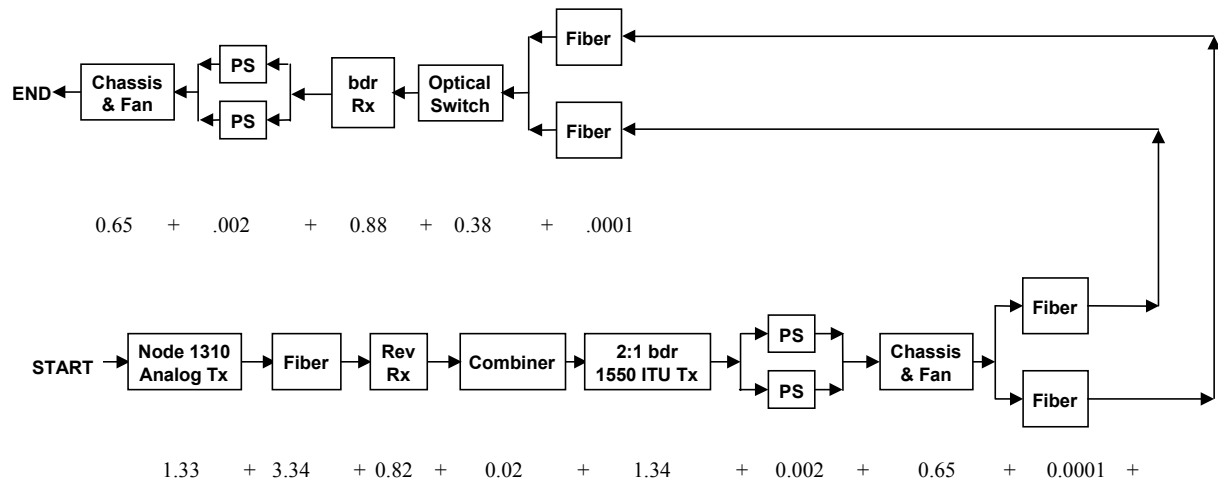
These RT locations are actually small, non-temperature conditioned cabinets that can be cost-effectively deployed on easements or on poles. This network expands the reach of the entire HFC plant, cost-effectively reducing Hub and Headend facility requirements.

Despite the incremental electronics and increased reach (85km v. 60km), the RT architecture is calculated to improve system availability by 10%. The key difference that enables this improvement is that the fiber links are better protected in the RT network.

Practically, the RT architecture cannot be deployed without the optical reach associated with bdr. From a reliability standpoint, the key issue is whether the reduction associated with the incremental electronics in the RT approach is outweighed by the shortening of the

As before, the complete set of outage time estimates for the RT architecture is shown in the table 15 following.

Figure 14 - RT Architecture Reliability Block Diagram



Total Annual Outage Time Per Subscriber = 6.74 minutes

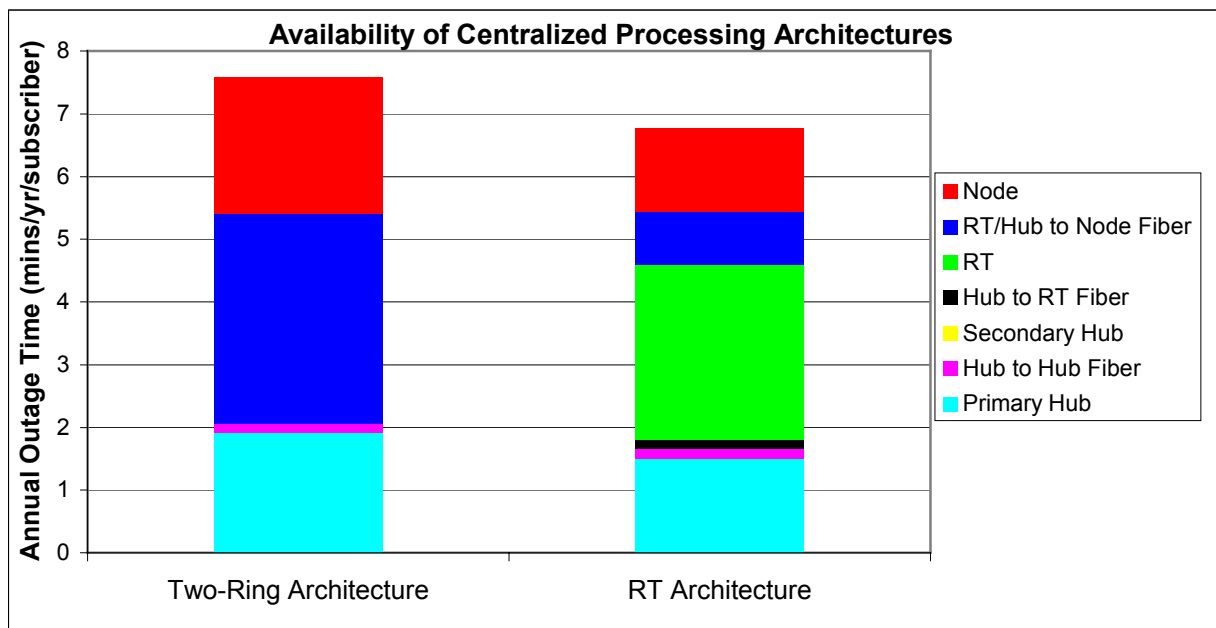
Figure 15 – Table of Outage Time Estimates

Location	Description	SA Part#	Type Redun	PCP - Predicted MTBF (hrs, M1C1 Q1 1)	Enhanced MTBF (hrs.)	MTTR (hrs.)	MTTC (hrs.)	Availability	Average Annual Outage Time (Min.)
Digital Reverse Path Traditional Two-Ring Architecture									
Primary Hub	ASSY MOD 14BIT PRISMA RCVR MOD UI F716157		2	139,225	1,489,145	24	2.50	0.99999832	0.8830
Primary Hub	Optical Switch P2 - SA Conn 714470		2	326,456	3,491,759	24	2.50	0.99999928	0.3766
Primary Hub	P2 Chassis & Fan 594300		2	189,449	2,026,338	24	2.50	0.99999877	0.6489
Primary Hub	Prisma II Chassis PS 589254		2		833,333,333	24	2.50	1.00000000	0.0016
Fiber	40 KM fiber NA		1/4+		314,941	4.00	NA	0.99998730	6.6800
								Primary Route Availability = 0.99998730	6.6800
								Availability with Redundancy = 1.00000000	0.0001
Secondary Hub	NA								
Fiber	20 KM fiber NA		1		629,882	4.00	NA	0.99999365	3.3400
Node	MOD ASSY 6940 BDR 2:1 DIGITAL MOD UI F712892		2	119,464	1,277,782	24	2.50	0.99999804	1.0291
Node	MOD ASSY 1560.61NM 6940 BDR XMTR 2.5Gbps 13312		2	106,408	1,138,135	24	2.50	0.99999780	1.1554
	Digital Reverse Path Total							0.99998566	7.4347
Digital Reverse Path Remote Terminal Architecture									
Primary Hub	P2-BDR-RP-2R 738959		2	169,043	1,808,077	24	2.50	0.99999862	0.7273
Primary Hub	Optical Switch P2 - SA Conn 714470		2	326,456	3,491,759	24	2.50	0.99999928	0.3766
Primary Hub	P2 Chassis & Fan 594300		2	189,449	2,026,338	24	2.50	0.99999877	0.6489
Primary Hub	Prisma II Chassis PS 589254		2		833,333,333	24	2.50	1.00000000	0.0016
Fiber	40 KM fiber NA		1/4+		314,941	4.00	NA	0.99998730	6.6800
Secondary Hub	NA								
Fiber	40 KM fiber NA		1/4+		314,941	4.00	NA	0.99998730	6.6800
Fiber								Primary Route Availability = 0.99997460	13.3600
Fiber								Availability with Redundancy = 1.00000000	0.0003
Remote Terminal	P2-BDR-TP-2R 738961		2	92,017	984,210	24	2.50	0.99999746	1.3361
Remote Terminal	P2 Chassis & Fan 594300		2	189,449	2,026,338	24	2.50	0.99999877	0.6489
Remote Terminal	Prisma II Chassis PS 589254		2		833,333,333	24	2.50	1.00000000	0.0016
Remote Terminal	Dual Reverse Rx P2 - Video SA Conn 716480		2	150,234	1,606,896	24	2.50	0.99999844	0.8183
Remote Terminal	8:1 SPLITTER/COMBINER FF 591816		2	6,830,600	73,059,805	24	2.50	0.99999997	0.0180
Fiber	5 KM fiber NA		1		2,519,526	4.00	NA	0.99999844	0.8350
Node	ASSY MFLX FP XMTR SCA 717904		2	92,563	990,050	24	2.50	0.99999747	1.3282
	Digital Reverse Path Total							0.99998718	6.7409

The graph in Figure 16. best depicts the calculated outage time per subscriber differences between the RT and more traditional Two-Ring Centralized Processing

Architecture. Readily apparent is the significant variation associated with the fiber links (approximately 2.5 minutes), which more than offsets the addition of the RT equipment.

Figure 16



Summary

The conversion from analog to digital technology has taken place across a wide spectrum of the products and services we touch daily – telecommunications networks, data communications networks, consumer electronics, household appliances, automobiles, etc. Baseband Digital Reverse represents the first step in the inevitable evolution of our HFC broadband networks to “all-digital”.

Today, bdr technology enables the use of new products and network architectures that are cost-effective, high performance, and reliable. The technology supports the performance required to deliver the higher order modulations schemes being adopted for data and voice services. The products demonstrate greater stability and can deliver more bandwidth per fiber than analog alternatives. The long optical reach capabilities and DWDM compatibility enable the use design and implementation of more cost-effective, highly reliable network architectures.

New applications for the technology are being investigated that will further improve cost, performance, and reliability of the fiber intensive networks operators will be deploying in the future.

References:

1. Andy Drexler, “Comparison of Node-to-Hub Availability for Analog and Digital Reverse Path Architectures”; December 7, 2000; Scientific-Atlanta Quality Assurance Project Document.
2. Andy Drexler, “Comparison of Fiber Deeper Remote Terminal Architecture Availability for Analog and Digital Reverse Path Approaches”; March 6, 2001; Scientific-Atlanta Quality Assurance Project Document.

SCALABLE BROADBAND ARCHITECTURES

Michael Field, System Architect
Ran Oz, Chief Technology Officer
BigBand Networks, Inc.

Abstract

More choice is a proven winner in cable television. In recent years, cable operators have made strides in providing more channels by extending the useful frequency range, and through digital compression. Powerful competitors such as direct broadcast and digital subscriber line service to the Internet continue to push for still more choice.

How can the cable operator continue to provide more choice to the customer while maximizing the returns on their investment in existing plant? One compelling way is by taking a page out of the data communications networking book and using switching technology to deliver the right mixture of services to each of the fiber nodes of their network.

This paper explores the way that switching technology allows the introduction of advanced services and thousands of new television channels.

THE TWO TIER ARCHITECTURE Headends and Hubs

Many cable systems consist of a central headend, in which satellite and off-air broadcasts are received, and hub sites from which fibers are distributed to neighborhood nodes.

Figure 1 shows this arrangement. The terminology 'Service Groups' is used to refer to a single fiber node or a collection of fiber nodes that serve a group of customers. Some operators refer to the Service Group as a 'Forward Carrier Path.'

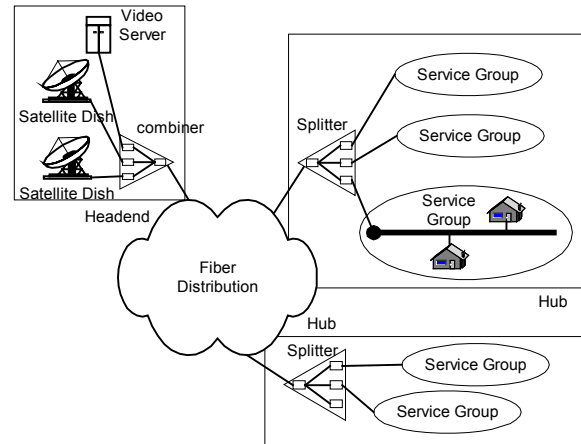


Figure 1 - Headend and Passive Hub Sites

This system works well for traditional broadcast video distribution as the signal can be split for distribution over the most economical means, optical fiber, to different neighborhoods and then finally delivered over coaxial cable to the home. Several variations of this type of network exist that increase the resilience of the fiber distribution against single component failure.

Expanding content availability strains the traditional passive broadcast model. In the same way that Ethernet Switching has increased the effective capacity of data networks, the capacity of broadband networks can be increased by replacing passive splitting elements with switching elements. This change allows the operator to customize the services offered and make maximum use of what proves to be the most limited, and costly, element of the system: the fiber-to-coax conversion and the final run of coax to the home.

Figure 2 shows the replacement of passive splitters with switching elements. With the

introduction of multimedia switching, the ability to provide locally inserted services, and more choice, becomes possible.

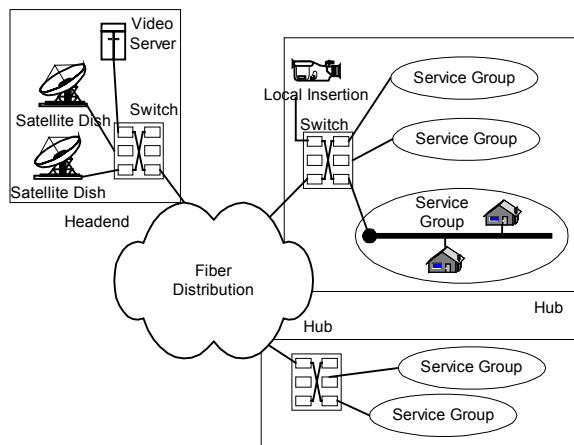


Figure 2 - Headend and Active Hub Sites

Tier One – The Headend

The headend contains facilities that benefit from centralization. Satellite reception, large capacity video servers and access to the Internet backbone are best located at this point in the network. The first tier of switching at this point allows customization of the services sent to each hub.

For smaller communities, direct switching from the headend to fiber nodes serving neighborhoods provides enough capacity per subscriber to support each subscriber adequately.

Tier Two - The Hub

A hub generally serves a geographic area. If sufficient demand for geographic diversity exists, the hub is the ideal site for the introduction of localized services through switching, for example a Video-on-Demand server for “Top 10” hits or culturally targeted programming. This not only provides more choice to the subscriber, but also has been shown to drive higher buy rates.

Fiber distribution to the hub can be implemented in two ways, with pre-bundled multi-program transport multiplexes or as large collection of single-program streams. If the first option is chosen the complexity at the headend is increased but the job of the switch at the hub is simplified, conversely, the opposite is true of the later option.

Switching at the hub provides the next great expansion in available bandwidth. This bandwidth can be utilized in simple ways, for example by providing different channel line-ups for different areas, insertion of ATVEF data, other interactive TV content or for narrowcast services such as Video-on-Demand. It can also be used in more innovative ways, described later in this paper.

Using Service Groups to Extend Bandwidth

The Service Group represents a collection of homes that share bandwidth. The second tier switching element shown in Figure 2 customizes program offerings and data delivered to each service group so that only the services required in that group occupy bandwidth.

This process is analogous to the actions of an Ethernet switch that improves performance on each leg of the network by switching only the packets required by that leg of the network.

Although not highlighted in this paper the devices that enable switching to the appropriate service group also enable insertion of services customized to that group, such as regionalized ad-insertion, VoD and in-band data delivery.

Implementation

Moving from today’s passive transmission systems to switched systems will be an

evolutionary process. As switching progresses deeper and deeper into the network, consideration must be given to existing network infrastructure, investment in set-top boxes and the scalability of switching devices. Fortunately, like the ever-growing branches of a fractal, switching systems can grow gradually to meet these needs.

To think of this another way – deployment of distributed switching in the network can be performed incrementally, the capital costs of such a deployment therefore rise in small steps. In contrast, deployment of central office based services requires a single giant step that represents a much larger risk.

Figure 3 represents current equipment arrangements at a digital headend. Various services are combined prior to distribution to fiber nodes.

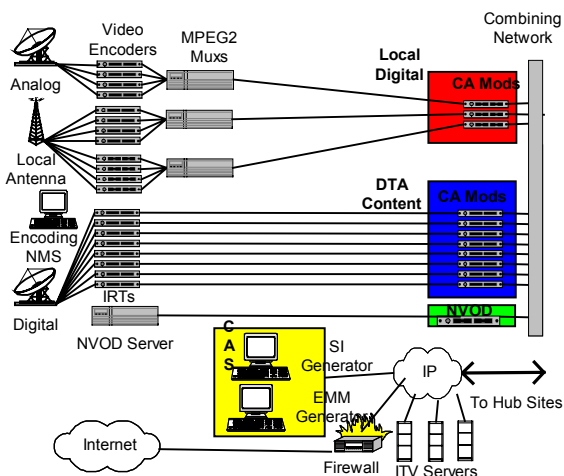


Figure 3 – Current Digital Headend

This picture changes when a switching device can prepare unique variations of services for each Service Group. This is illustrated in Figure 4.

Finally, when distribution of processing power is required at the hub site the picture is again transformed to Figure 5.

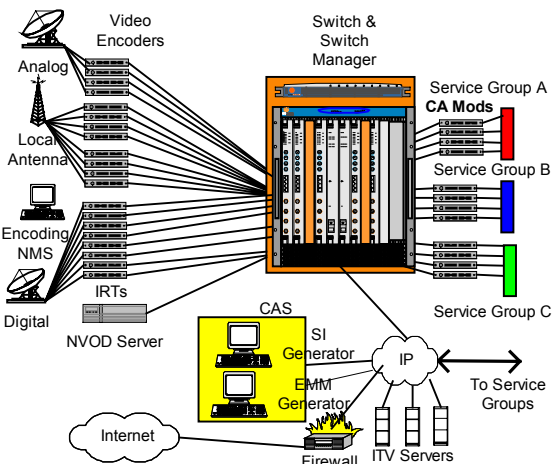


Figure 4 – Switching to Service Groups

Distribution of processing capacity to the hub:

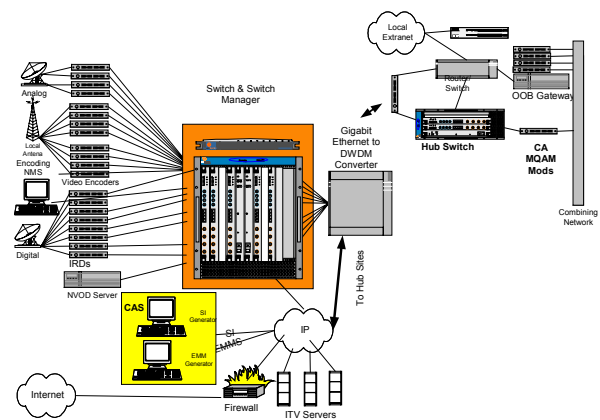


Figure 5 – Switching at the Hub

Switching Technology

Several choices of switching technology are available for use at tier 1 or tier 2. These include:

1. ATM switching using AAL5 encapsulation.
2. Ethernet switching using one of several physical Ethernet interfaces (e.g. 100baseT, gigabit optical.)
3. MPEG-2 transport switching with industry standard DHEI and ASI interfaces.
4. SONET.

SONET and ATM switching are widely used in telephony and in some existing cable plant. However, switch costs and flexibility will increasingly favor gigabit Ethernet in coming years for fiber distribution. Switching of synchronous MPEG-2 transport and interfaces to existing MPEG-2 equipment (such as IRTs) play an important role in the headend and hub, such interfaces must continue to be supported.

In addition to the star topology described in Figure 2, cable networks have been implemented using a ring connecting each hub. The ring somewhat reduces the need to switch data at tier 1. It is, however, still beneficial to combine channels into statistical multiplexes at this level, perhaps even carrying certain channels more than once in order to reduce the switching load at the hub. Clearly this is not theoretically needed but is an example of an optimization that can be performed for economical deployment.

Video Processing

Unlike many data streams, audio and video information cannot be flow controlled to suit the recipient of the data. Producing a mixture of channels for each neighborhood can be made significantly more economical if video processing techniques are used in delivering the content. Producing a statistical multiplex of variable bit-rate video is a form of flow control that improves the usage of a transmission channel.

NICHE PROGRAMMING

Because the final stretch of cable to the home has a finite capacity the operator has been placed in the often-awkward position of choosing among available sources the television channels that will best serve their community. As such, individual subscriber choice is limited because of the need to

consider the whole community or comply with must-carry regulations. Expanding the numbers of channels available and providing the flexibility to switch only those requested to the neighborhood enables the operator to be one of the 'good guys' with little risk of long term commitment of transmission network bandwidth.

Broadcast-on-Demand Switching

Broadcast-on-Demand, the ability to join a channel or program in-progress, and only burden the local cable segment with the cost of carriage for that channel, is made possible with three technologies:

1. Economical high-speed multimedia switching that directs the needed channel to the correct branch of the transmission network. Economical switching necessitates the use of statistical multiplexing and video rate-shaping.
2. Multicasting to enable a single copy of a channel to be shared on single branch of the transmission network.
3. User directed channel selection. This could be dynamic, e.g. the set-top box signals channel changes, or less dynamic, e.g. conventional customer subscription.

In order to implement Broadcast-on-Demand the set-top box signals the desired channel upstream to a Switch Manager. The Switch Manager uses information about the desired channel and the user's Service Group to signal the switch at the headend (tier 1) for the desired channel to be transmitted to the appropriate hub. The Switch Manager then signals the switch at the hub (tier 2) to transmit the channel to the user's service group. Subsequently the set-top box tunes to the appropriate frequency and program.

If the desired channel is already being switched to the hub or the appropriate service

group, no additional switching actions are required and the set-top can join a multicast of the channel. The Switch Manager maintains a count of users for the channel so that it can remove the channel from a specific service group when all its users have dropped out.

Channel Surfing Switched Programming

Ideally the user should be able to surf through all available programming in the manner to which they have grown accustomed since the introduction of the TV remote control. Digital television has already affected the users ability to do this due to the long acquisition time required to capture the first frame of a program. The user has responded by flipping through channels, reading only the information bar describing the programming and finally alighting on a program of interest. Given enough transaction processing capacity a Broadcast-on-Demand solution need only add modestly to this wait for the first frame of video to appear.

Several related options exist that the public may embrace; each of these is enabled by switching:

1. Channel choice through the Electronic Program Guide.
2. A surfing channel that cycles quickly through available channels.
3. A surfing mosaic, giving the user instant access to several channels simultaneously.
4. Conventional channel surfing

The Electronic Program Guide has long been a focus for choosing available programs. Some users have embraced the guide, but others continue to click-click through programs because it is easier than using the guide, and possibly because of the dullness of seeing the choices printed.

A surfing channel has the advantage of eliminating the acquisition time between digital channels – a disadvantage is that most surfers click quite quickly.

A surfing mosaic again eliminates the acquisition time between digital channels. Several channels are displayed in a matrix so that all are simultaneously visible. An additional enhancement would allow the user to surf the sound on each channel as it can be processed more quickly.

Finally, the same model the user currently uses – simply pausing on a channel – note: when the user pauses on this channel the channel change at the tier 2 switch occurs.

Video-on-Demand

It should be noted that Video-on-Demand programming supplied by file servers shared many of the same switching problems that are solved by switching for Broadcast-on-Demand. It would be useful if the Broadcast-and Video-on-Demand interfaces were shared. This would present the customer with a full range of options in a unified manner and allow them to choose among them.

Advanced Services

Using the switched architecture described above more advanced services such as in-band data and ad-insertion as possible. These can be integrated with broadcast content. The localization of such switching allows powerful opportunities for targeting.

Network Usage

Carrying the right mixture of programming to each neighborhood implies keeping track of the channels in use, adding to them when requested and removing them when no-one is watching. The aggregate time that a niche

program is in use is surely a useful statistic for the transmission network planner and network advertiser.

Conclusion

The segmentation of the broadcast cable television network into units that serve a smaller group of customers makes possible an increase in the number and quality of services delivered. Switching devices enable the appropriate mixture of services to be tailored to each segment of the network, for economical deployment these switches need to include features that accommodate legacy equipment such as existing MPEG-2 set-top boxes and multi-QAM modulators. Extending switching concepts to broadcast services extends these advantages, with potential to remove bounds on what content can be offered to subscribers and to expand functionality.

BigBand Networks has implemented a Broadband Multimedia-Service Router, that operates using a hybrid Ethernet/ IP / MPEG 2 transport technology. Powerful video processing allows efficient use of standard MPEG-2 transport streams. The switching unit is complemented by a Service Manager, using a variety of protocols, including industry standards such as DSM-CC.

References

1. Optimal Network Architectures for On-Demand Services – SCTE Emerging Technologies Proceedings.
2. Modern Cable Television Technology – Video, Voice, and Data Communications – Ciciora, Farmer, Large.
3. OpenCable Architecture – Michael Adams.
4. Unified Headend Technical Management of Digital Services – Ran Oz, Amir Bassan-Eskenazi, David Large.

SIMULATION OF VIDEO-ON-DEMAND TRAFFIC¹

Mark Cronshaw, Ph.D. & Victoria Okeson
AT&T Broadband Labs, Westminster, CO

Abstract

This paper describes a discrete-event simulation model of VOD traffic in a system with a server and hierarchical network. The model predicts the magnitude and location of blocking as a function of the demand for VOD sessions and their durations. We found that an Erlang-B model may overestimate the blocking probability, since it is based on a statistical equilibrium that may not be reached with time-varying VOD traffic.

The model provides an inexpensive way to explore various network configurations and modulation schemes. As such, it can be very useful for VOD capacity planning.

INTRODUCTION

Interactive cable traffic such as video-on-demand (VOD) content presents fascinating capacity planning issues. For example, VOD systems must be sized so that there is only a small probability of a subscriber being unable to initiate a session. Demand for sessions varies by time of day and day of week. Also, the duration of a VOD session depends on the content selected, the time spent selecting the content, and on the use of

VCR-like features such as pause, rewind and fast-forward.

In some respects network traffic associated with video on demand is similar to telephony traffic. Requests for new sessions arrive randomly over time, and the duration of each session is random. Furthermore, the average arrival rate of new requests is not constant over time, but rather there is a peak period during any given day, and this peak will vary on different days. For example, the peak demand for VOD will probably be on a cold and rainy Friday night. However, there is an important difference in the nature of the traffic: the average length of a VOD session is long relative to the duration of the peak period. In contrast, the duration of a typical phone call is short relative to the duration of the peak calling period.

Figure 1 is a stylized representation of VOD traffic. Suppose that the peak period lasts for 3 hours, and that movies are only requested during the peak period, at a uniform rate of 10 per hour. Suppose also, that each VOD session lasts for exactly 2 hours. The solid curve in the figure shows the arrival rate of requests. The dashed curve shows the rate at which movies end: it is the same as the arrival rate curve but delayed by 2 hours.

¹ The authors would like to acknowledge many indispensable conversations with Jim Want and Doug Ike, and to thank Rodger Woock for his review of this paper.

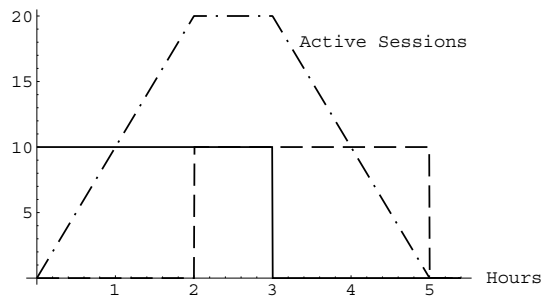


Figure 1 Stylized Video on Demand Traffic

The dash-dotted curve shows the number of sessions that are active, assuming that there are no capacity constraints. During the first two hours the number of active sessions grows at a rate of 10 per hour, reaching a maximum of 20 sessions. Beginning at hour 2, movies finish that had started earlier. Simultaneously, new requests for movies arrive, so there is no change in the number of active sessions between hour 2 and hour 3. The peak period ends at hour 3, after which there is no demand for new sessions. From hour 3 to 5 movies finish which had started previously, and the number of active sessions falls to zero.

By way of contrast, Figure 2 shows the number of active sessions for stylized telephony traffic assuming that the peak period runs for one hour with a uniform arrival rate of 400 calls per hour, and that each phone call lasts for exactly 3 minutes.

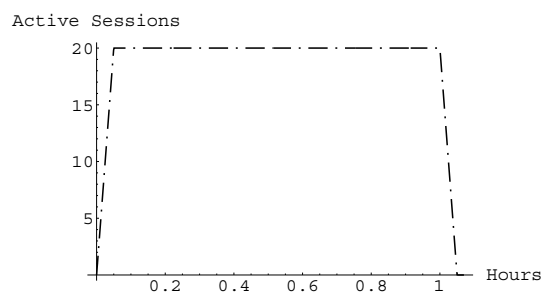


Figure 2 Stylized Telephony Traffic

The number of active sessions (phone calls) over time has a similar shape to that for the VOD traffic, but there is an important qualitative difference. The number of active phone calls is at a sustained maximum level for almost all of the peak period. In contrast, the maximum number of active VOD sessions is at its maximum level for only a small part of the duration of the peak period. This is a direct consequence of the relatively long length of a VOD session compared to the duration of peak demand for VOD.

This qualitative difference has important implications for modeling VOD traffic. The Erlang-B model that predicts blocking of telephony traffic is based on a stochastic equilibrium, which corresponds roughly to the long sustained period of maximum sessions in Figure 2. Under typical conditions, VOD traffic may not reach such a stochastic equilibrium, and hence the Erlang-B model [Wall] may not provide a good approximation of blocking. In fact, if the network is initially empty, then no blocking will occur at the beginning of the peak period. Blocking will only occur after the network has filled to a level where a capacity constraint is reached.

This paper describes a numerical simulation model of VOD traffic that predicts blocking. It shows that

- blocking increases rapidly when the demand for VOD sessions exceeds a threshold,
- blocking is sensitive to the average duration of a VOD session, and
- the Erlang-B model may overestimate blocking.

SIMULATION FRAMEWORK

Network configuration

We modeled a network architecture with a central VOD server that supports several hubs. Video streams are carried from the VOD server to each hub on fiber. Each hub delivers VOD traffic modulated onto RF channels to several nodes. The RF channels are shared among groups of three nodes, called supernodes. Thus there are three potential locations for blocking: at the server, in the fiber between the server and the hub, or in the RF.

The results in this paper are for a network with 1 server, 12 hubs, 217 supernodes and 4 RF channels per supernode. The number of supernodes per hub ranged from 6 to 27. The modulation of each RF channel was 64 QAM, providing 8 digital streams at 3.375 Mbps.

Characteristics of VOD sessions

For the purposes of traffic modeling the relevant characteristic of a VOD session is its duration.² This is determined by

- the length of the content,
- the time used while selecting the content,
- the proportion of sessions which are terminated before the end of the content, and
- the time used for VCR-like features such as pause, fast-forward and rewind.

We modeled two distinct types of sessions: long and short, corresponding

to feature length films and shorter content such as children's programs. We assumed that the durations of each type of session are normally distributed with means of 120 and 55 minutes respectively, and standard deviations of 25 and 5 minutes respectively.

Demand for VOD sessions

The demand for each type of session varies by day and by time of day. Our simulations cover a single day, which for capacity planning purposes should be the busiest day anticipated. A natural unit for demand is the number of buys per subscriber per month. However, for the purposes of simulation, it is necessary to specify how these buys are distributed over time. We assumed that session requests are distributed according to a Poisson distribution, and that the demand is the same for each supernode. Demand is specified in terms of an average arrival rate of session requests. Different demand can be specified for each type of session.

Demand for each type varies by time of day. For each type the day was divided into four phases. The average arrival rate in each phase was specified as a percentage of the rate during the peak phase, as shown in Table 1. The table shows time based on a 24 hour clock. We assumed that the peak rate was the same for both long and short sessions (although the peak for each occurs at different times, as shown on Table 1.)

² We only considered traffic associated with VOD content. There is also out-of-band traffic associated with control of a VOD session, such as ordering, pause, fast-forward and rewind.

	Early		Pre-peak		Peak		Late	
Type	Time	Demand	Time	Demand	Time	Demand	Time	Demand
Long	0300-1700	5%	1700-1900	30%	1900-2200	100%	2200-0300	5%
Short	0300-1400	5%	1400-1600	30%	1600-1900	100%	1900-0300	5%

Table 1 Average arrival rates (as a percentage of peak rate)

We considered peak average arrival rates requests (for each type) between 2 and 5 per hour per node. Based on the time varying demand, a peak rate of 2 requests per hour per node is equivalent to 18.2 requests per node per day. For a node with 500 homes passed and 11% digital penetration, this means that the average number of session requests per digital home in that day is 0.33. If the total demand over all of the other days in the week were $\frac{1}{2}$ of that on this busy day, then the monthly buy rate would be $4 \times 1.5 \times 0.33 = 2$ buys per digital household, based on a four week month. The buy rate scales linearly with the peak rate. So a peak rate of 4 requests per hour per node corresponds to 4 buys per month given these assumptions.

Network performance

There are two aspects of blocking that are of interest

- a customer perspective: the percentage of session requests which are not granted, i.e., blocked, because of capacity constraints, and
- a network perspective: the location of blocks in the network, which can be at the server, in the fiber between the server and the hub, or in the RF.

Since demand varies over time, the blocking probability also varies over time. The simulations generated

blocking probabilities for each hour. The blocking probabilities reported in this paper are for the hour with the greatest blocking. For most of the runs this was from 9 to 10 pm. For some of the runs it was between 8 and 9 pm. Based on Table 1, the highest average arrival rates occurred from 5 to 7 pm. The highest blocking occurs later than the period of peak demand, since the network is relatively empty at the beginning of the peak request period, but fills as session requests are granted. As shown on Figure 1, it takes some time for the number of active sessions to grow to a point where a capacity limit is reached.

MODELING APPROACH

The model was implemented using a discrete event Monte Carlo simulation written in Microsoft Excel and Visual Basic for Applications. The user has the option to enter up to five network configurations and three demand scenarios. Each unique network configuration – demand scenario pair is called a case, and the user must select the specific cases to study. The model treats each case independently – this mechanism simply allows for more efficient processing of multiple cases.

The model is designed to simulate a peak day in the system. Movie demands are defined for a 24-hour day, starting at 3 am. The model considers 3 am to 7 am

to be a warm-up period to initialize the network, so no statistics are collected during this time. At 7 am, the model begins collecting statistics for the next 20 hours. Two types of statistics are collected: blocked requests and active sessions.

Because this is a Monte Carlo simulation, multiple trials must be run for each case. At the beginning of each trial, movie requests are generated according to a Poisson process for each type of movie at each node in the network. These requests are then sorted by arrival time and processed in turn. Processing a request has three steps: clearing out completed movies, assessing the network for capacity and reacting to the result of the assessment. First, the entire network is checked to see if any movies have completed since the last movie request was processed. If so, the capacity counters in each network element are updated to reflect the newly released capacity. Next, the supernode at which the movie arrives is identified. This supernode, the fiber that serves it, and the VOD server are all checked for available capacity. If sufficient capacity exists in all of these network elements, then the movie request is granted and the duration of the movie is obtained according to a predefined statistical distribution. The movie is recorded in the network and capacity counters are updated. If there is not enough capacity, then the location of the block is recorded. This process repeats until all of the requests are handled. The active sessions data is recorded by evaluating the state of the network at user-specified intervals throughout the course of the trial.

Once all of the trials are completed, summary reports, which include the averages over all of the trials as well as

individual trial results, are output for later analysis.³

Validation

Validating the model is essential to insure that it accurately represents the system under consideration. For this model, an analytic benchmark is available for validation. A VOD system resembles an M/M/m/m queue, which has multiple servers that block requests if no server is idle when the requests arrive. Applying this to the VOD simulation model is something of a generalization, since a VOD system is a hierarchical network. However, for the network we modeled, the vast majority of the blocking occurs at the supernode level, so it is a reasonable approximation. For an M/M/m/m queue, the blocking probability (i.e. likelihood the request is lost) is simply an Erlang-B function [Tanner]. We compared the analytic benchmark to the VOD simulation model for a system with only one movie type. Table 2 illustrates the comparison between the analytical blocking and the VOD simulation model blocking for three different cases, varying only the movie request arrival rates. The arrival rates were kept constant over time in the simulations to satisfy the stochastic equilibrium assumption implicit in the Erlang-B model. Note that our simulation model shows excellent agreement with the analytic results.

³ Since demand is random, the actual highest level of blocking may be higher or lower than the average. There are other criteria besides average blocking that might be used for network design, such as a 5% chance that more than 1% of session requests are blocked.

Average session request arrival rate per supernode per hour	Analytical blocking	VOD simulation model blocking	% of blocking occurring at supernodes
15.9	4.8%	5.0%	100%
17.7	8.8%	9.2%	100%
19.2	12.9%	13.4%	99.8%

Table 2 Model Validation

MODELING RESULTS

Table 3 shows how the blocking probability varies with the demand for VOD sessions. With time varying demand, there is no blocking if the peak average arrival rate is 2 requests per type per node per hour (or lower). On average, 3% of requests are blocked if the peak rate is 4 requests per type per node per hour.⁴ But if the peak rate rises to 5, then blocking reaches an unacceptable level of 12.7% of session requests.

Peak avg. arrival rate per type per node per hour	Demand	
	Time varying	Steady
2	0.0%	0.05%
3	0.15%	4.2%
4	3.0%	18.7%
5	12.7%	35.2%

Table 3 Blocking probabilities (highest hour)

The table also shows the blocking probabilities if demand is constant at the peak rate over time, rather than time-

⁴ Blocked requests were not queued in the simulations. They were simply not granted. In reality a blocked request would probably lead a subscriber to submit another request immediately. This behavior was not modeled.

varying. The blocking is significantly higher with steady demand. This shows that the Erlang-B model can overestimate blocking if the demand for VOD sessions varies over time.

Table 4 shows the location of blocking in the network, expressed as a percentage of all blocks. All blocking occurred in the RF, except for a small amount of blocking at the server when the peak average arrival rate was 5.⁵

Peak avg. arrival rate per type per node per hour	Location		
	Server	Fiber	RF
2	none		
3	0%	0%	100%
4	0%	0%	100%
5	0.8%	0%	99.2%

Table 4 Blocking locations

The blocking probability is sensitive to the average duration of a VOD session. If the average duration of the long

⁵ It is possible for there to be insufficient capacity simultaneously in several parts of the network. We attributed blocking to the lowest level of the network where it occurred. For example, if both the RF and the server were full when a new request arrived, then the block would be attributed to the RF and not to the server.

sessions is reduced from 120 to 110 minutes, and that of the short sessions from 55 to 50 minutes, then blocking is approximately halved. With time varying demand and peak rates of 4 and 5, the blocking probabilities are 1.3% and 7.4% respectively, compared to the values of 3.0% and 12.7% in Table 3. A drop in blocking would be expected due to the reduction in offered load. Yet the magnitude of the drop is large. This sensitivity demonstrates the potential benefit from reducing the average VOD session length, perhaps by taking steps to reduce session time involved in movie selections.

We also did simulations to explore a lower standard deviation of the session duration. Such reduction had only a small impact on the blocking probability.

CONCLUSIONS

Simulation is a useful tool to assist in capacity sizing. It is flexible, which makes it possible to evaluate the anticipated performance of many possible network configurations, before making significant investments.

The discipline of building a simulation model brought to light many crucial operational issues, such as

- the importance of the average duration of a VOD session, and
- the time-varying nature of demand and blocking.

It would be highly desirable to collect network data to calibrate the model. The model predictions do match the analytical Erlang-B model. But the Erlang-B model is not valid for time-varying traffic on an actual hierarchical network. Calibration with actual VOD session traffic would increase confidence in the predictions of the simulation model. After such calibration, the model can generate significant money-saving insights on how to achieve a balance between the costs of adding VOD capacity and the costs of having subscribers frustrated by blocked VOD requests.

REFERENCES

Tanner, Mike. *Practical Queueing Analysis*. McGraw-Hill, Berkshire, 1995.

Wall, Bill. "Traffic Engineering for Video-On-Demand Systems", Scientific Atlanta, 2000.

Contact info

Mark Cronshaw

mbcronshaw@broadband.att.com

Victoria Okeson

vokeson@broadband.att.com

10355 Westmoor Dr., Suite 100
Westminster, CO 80021

SYNCHRONIZING DEEP FIBER BASEBAND ACCESS NETWORK DESIGN WITH TRADITIONAL HFC INFRASTRUCTURE

Donald Sorenson
Scientific-Atlanta, Inc.

Abstract

Increasingly optical network technologies are being evaluated for their suitability as a residential access network platform. Clearly optical networking holds promise for the distant future as networked applications evolve to demand greater transport capabilities. However, today and for years to come HFC based access systems with their low cost structures and evolving performances will remain the dominant access network of choice for delivery of interactive multimedia services in the majority of the served residential markets. In a contemporary context it is quite likely that an all-optical access platform could be best utilized as a strategic tool tailored to delivering high-value business class services to 10% or less of a residential serving area. A low first cost optical access network may be an ideal strategic/offensive overlay to an existing or new HFC network.

This paper explores the key applications engineering issues associated with such an overlay and proposes a methodology for synchronizing key optical access and HFC network elements. The paper concludes with a detailed analysis of how optical split ratios and cable sheath fiber counts can impact plant first costs. This work is focused on the engineering issues associated with layer 1, physical layer, of the 7-layer OSI network model. Commentary on upper layer requirements and protocols are limited to issues that impact the logistics of an overlay deployment and subscriber provisioning.

HFC AND FTTx ARCHITECTURES

The overlay of contemporary deep fiber architecture on an existing or new HFC plant is facilitated by effectively coordinating the location and functionality of key network elements. A brief review of each access method's legacy and contemporary structural fundamentals will serve as a useful point of reference as various design and cost sensitivities are explored later in the analysis to follow.

Hybrid Fiber Coax (HFC)

Routinely the architectural details of HFC plants are varied to meet a host of constraints associated with factors such as geography, home density, up-grade logistics, costs as well as many other details. However fundamentally most plants adopt a variant of a ring(s) – star(s) – bus configuration, as illustrated in Figure #1. Typically one or more fiber rings interconnect primary headend and hub facilities, in some cases secondary rings link headends with secondary hubs or **optical termination nodes** (OTN's) that commonly serve tens of thousands of subscribers. In turn fiber nodes are typically linked to hubs, OTN's or headends in a star configuration.

Ring and meshed configurations offer physical path diversity and associated enhancements in plant reliability. As high performance service and business models stimulate the need to extend fiber deeper into the plant a greater use of rings is anticipated at deeper levels. Within this trend network

functionality is also on the move, network process functions are migrating to the edge of the plant where they can be better tailored and incrementally scaled to meet subscriber and new business demands. A third ring extending from a secondary hub through a series of smaller serving areas to remote star locations at the one to four thousand homes passed level effectively enables the deployment of advanced video and cable data services [1] and, is a logistical requirement for the deployment of an all-fiber access overlay, as will be shown later in this paper.

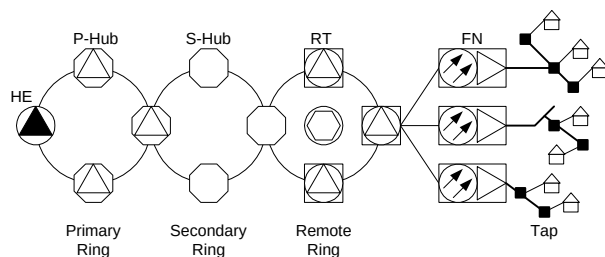


Figure #1

Optical Access Network Overlay

Various all-optical access networks are increasingly seen as:

- The enabler of high value services to a large underserved population of small business and home offices.
- A low risk infrastructure enhancement that has the potential to extend a plant's useful life well beyond that of copper based technologies.

A contemporary all-optical access architecture must be optimized to cost efficiently deliver advanced services to a small percentage of a wide subscriber market base being passed.

To insure the integrity of these services and an extended useful plant life the network must be reliable and easily scaled with increasing take rates without major upgrade or modification.

By contrast to most copper based networks all-optical access networks do not utilize active amplification or signal processing elements in the transmission path and are referred to as **passive optical networks** or PON's. Ideally the active elements at both the subscriber and service provider ends of the network are safely housed in conditioned facilities, and only passive components such as fiber cable and splicing devices suffer the environmental rigors of the outside plant, as shown in Figure #2.

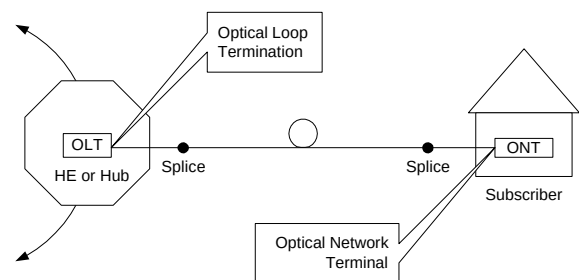


Figure #2

Rarely would a dedicated fiber connection between the service provider's internal network and a residential or small business subscriber be cost effective. To distribute costs a means of allowing multiple subscribers to access a network fiber is implemented through the use of multi-port optical splitter or coupler, Figure #3a.

In this instance each PON fiber becomes a shared media. Downstream traffic is broadcast to all of the subscribers served by a splitter, and upstream access is controlled via a suitable **media access control** (MAC) protocol, a number of which are competing for industry acceptance. For reasons relating to cost and operational logistics it is often not desirable to lump the splitting function into one physical location [2,3]. As shown in Figure #3b the splitting function can be divided between multiple locations along the fiber path. The product of the split ratios at

each location becomes the PON's overall return fiber to subscriber ratio.

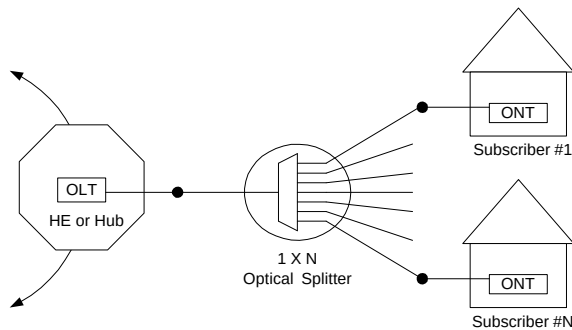


Figure #3a

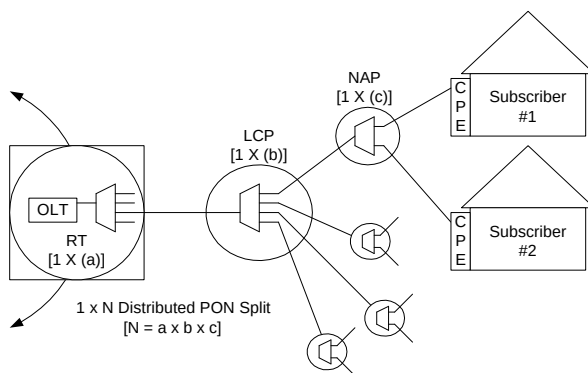


Figure #3b

The sizing and distribution of the PON splitter function is the principal tool used to synchronize an optical access network's elements with an existing or new HFC plant.

As previously mentioned, ideally a PON's active elements are housed in conditioned environments. However the realities of applying an optical network design to a wide serving area typically demand that active electronics be placed in the outside plant. At the subscriber side of the network the logistical realities of service activation place the **optical network terminal** [4] (ONT) outdoors on the exterior of the subscriber premise as **customer premise equipment** (CPE). At the network side construction costs and operational issues typically demand a high degree of fiber aggregation deep in the design. Without such measures fiber counts quickly become difficult to manage with

even the highest of PON split ratios. Network reliability and scalability suffer; typically it is not practical to provide path redundancy on one or a few PON fibers, however when hundreds of fibers collect on their way to a network facility the failure group size and recovery time from a fiber cut become unacceptable. Further, high fiber-count headend or hub runs are very likely to exceed existing fiber inventories and subsequently adversely impact plant first costs.

*The location of **optical loop termination** [4] (OLT) synchronizes quite well with the evolution of the third ring nodes or **remote terminals** [1] (RT's) in advanced HFC architectures.*

Commonly an OLT or RT location will service a 1000 to 4000 home area. Construction practices demand the consolidation of fiber cable sheaths as they route to the RT. These **local consolidation points** [2] (LCP's), Figure #3b, provide another physical location to easily distribute a portion of the network's total split ratio.

The proper allocation of PON split ratio to the LCP can be used to synchronize its physical location in the network with that of each physical HFC fiber node, thus providing a convenient fiber path for the HFC node and common physical point of equipment access.

Access to the PON fibers passing a subscriber premise is achieved in manner similar to an HFC coax tap. The principal difference being that the equivalent PON tap is fed by a dedicated PON fiber that is not shared by other up or downstream PON taps, Figure #3b. To logically clarify this difference the PON fiber tap point is referred to as a **network access point** [2] or NAP. Fundamentally a NAP is composed of a

suitably sized ground or aerial mounted fiber splice enclosure near a small group of subscribers. The optical service drop to the subscriber is terminated at the CPE using an optical splice or connector. At the NAP the drop will be spliced either directly on one of the PON fibers passing through the NAP or onto a port of an optical splitter located in the NAP.

The installed first cost and service activation cost of a PON are quite sensitive to the fraction of the total PON split ratio allocated in the NAP. In general a higher split ratio at the NAP yields a lower first cost of construction and slightly higher cost per subscriber served should the take rate at the NAP exceed 50% of its capacity.

Another challenging issue regarding PON architectures is that of remote active element powering. Historically remote network device power has been supplied via the metallic signaling media. Without such a media a PON must either include some form of copper based powering network or its remote actives must be locally or subscriber powered. Both alternatives have their advantages and difficulties, a selection between which will largely be determined by business and public policy issues.

The issue of powering remote PON equipment in an HFC overlay is facilitated by utilizing the powering signal commonly available from the underlying CATV distribution coax.

By including a light gage twisted pair with the subscriber's optical drop network power can be provided to the CPE using industry standard power passing tap devices. In those cases where spare power capacity is not available from the HFC network other twisted pair powering alternatives exist that can be implemented at the time of plant

construction at a modest percentage of total PON overlay cost.

NETWORK DESIGN CONSTRAINTS

The practical objective of a design is the implementation of a serviceable network that meets near and long term needs for the lowest possible cost. Compliance with this objective in legacy copper based technologies are strongly enforced by well understood and refined engineering practices. In part the scope of these practices must be expanded when new technologies such as PON are adopted. The following material outlines key engineering, construction and operational considerations that will act to constrain the design variables associated with implementing a PON overlay.

Serving as a reference point for the following discussion, Figure #4 is a simplified illustration of a PON overlaying an HFC fiber node with a physical size 256 **homes-passed** (HP). In turn the **physical node** (PN) is fed by a 1000 HP **logical node** (LN) or **remote terminal** (RT).

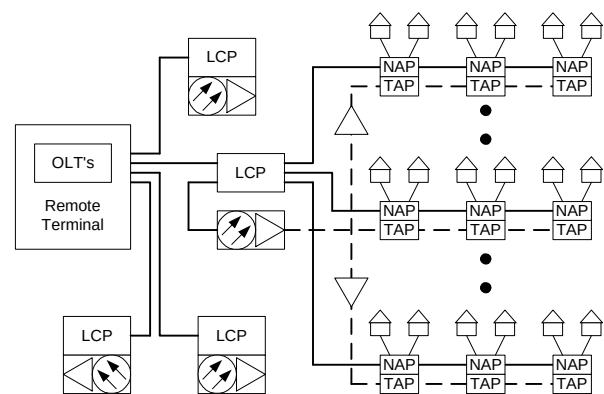


Figure #4

Network Engineering Considerations

Routinely more than one communications service provider serves a residential or commercial market area. This implies a reduced initial service activation rate as

market share is obtained. Accordingly new plant designs of any type are commonly engineered for a lowest possible first cost that enables the network to physically address the market area. This condition is particularly true with regard to PON deployment. It is anticipated that early take rates for PON based services will be quite low, less than 10% of HP. This condition is likely to remain until the value of service portfolios, exclusively deliverable via an optical network, increase and construction cost associated with PON decrease.

With a lowest possible first cost a PON overlay must address 100% of a market service area, and with minimal provisioning or manipulation the network can be made available any one of the subscribers passed.

In keeping with a low first cost and take-rate assumptions it is anticipated that PON split ratios will be under subscribed by as much as 50% and that initial RT data processing capacity will be oversubscribed by as much as 10:1. These two factors substantially lower first costs with regard to OLT ports, associated packet switching and fiber counts and/or advanced optical wavelength provisioning in the upper network transport rings. A 50% under subscription of a PON split ratio implies that less than 50% of a PON's available optical splitter ports will be terminated and in service at any one time. To enable this engineering option the selected MAC protocol must have the ability to dynamically allocate access to the PON over a specified range of possible subscriber counts. This capability greatly simplifies provisioning and record keeping while allowing the network to provide bandwidth to customers based upon the terms of their subscription agreements without regard to the number of subscribers terminated on a given PON. Conversely, and based upon optical budget limits - discussed later, under

subscription allows the network engineer to pass lit fiber by all of the homes in a given serving area without excessive OLT costs.

Selection of the proper MAC protocol and the planned under subscription of a multi-port PON allows the network to distribute lit fiber past 100% of the subscriber base without encumbering excessive OLT provisioning cost or degrading subscriber data throughput.

If the PON is to be under-subscribed it is advisable to take measures that help ensure the under subscription through the life of the network. This can be achieved by distributing a portion of the PON split ratio to the LCP. By doing so each of the PON's ports can be spread across multiple distribution branches or NAP locations that are not geographically adjacent. Routinely subscription rates in small clusters of homes passed can reach 100%, by spreading a PON's ports across multiple distribution branches the likelihood of 100% utilization is diminished.

The PON network must be capable of supporting the eventual migration of standard services delivered today via copper media. This requirement principally impacts optical budget limits that in turn drive PON split ratios and passive device performance limits. Large PON split ratios such as 1:32 demand optical **physical interface** (PHY) transmitter power levels of up to +20dbm. Elevated levels such as this are achieved at a cost premium and bandwidth penalty over more common PHY transmit levels of +10dbm or less.

In years to come the need to operate the PON at multi-gigabit speeds is almost certain. Large PON split ratios and the associated elevated PHY transmit power levels are likely to delay the point in time where an

existing network can be easily upgraded to multi-gigabit operation.

Thus PON splits such as 1:8 or 1:16 are attractive from the point of keeping CPE PHY-device costs down and ensuring early migration to higher data rates. It may also be advantageous to allocate some portion of the PON split ratio to the OLT side of the network. For example, a 4-port directional coupler can facilitate insertion of overlaying λ 's to support additional services in the future.

Construction Engineering

Minimizing construction first costs are key to the financial success of a PON overlay. Low first cost enables early economic business models that rely on the delivery of high value services to a selected few of the subscribers passed by the network. Unfortunately only a small number of PON networks have been deployed and related experiential knowledge is limited at best. The following points address construction issues impacted directly by PON architecture decisions.

Through the analysis to follow it has been found that PON construction costs are reasonably insensitive to actual fiber counts. However costs are particularly sensitive to the labor associated with fiber management, termination and optical splitters along with related outside plant enclosures and places to locate them. Again the size and distribution of PON split ratios can be optimized to minimize these initial costs.

Analysis has shown that poorly designed 1:32 or 1:16 PON can be significantly more expensive than an optical network based on a dedicated home-run fiber for every home passed!

PON fiber counts at a single location in the network can easily escalate to levels that are exceedingly difficult and costly to manage. Additionally, in the event of a fiber cut a large count represents a correspondingly large failure group size. For example; 200 - 1:16 PON fibers under subscribed by 50% still represents 1,600 high-value paying subscribers that will lose service when a pole falls or a cable is dug up. With cable restoration time in mind it is advisable to keep individual fiber sheath counts below 50 if possible and fiber aggregation point (LCP and RT) counts well below 150. This can be achieved by adjusting the amount of PON split ratio performed at the NAP and LCP.

Increasing NAP split ratios reduce fiber counts at the LCP and move the LCP into the network, which can be used to synchronize the LCP location with an existing or planned HFC fiber node. Correspondingly higher LCP split ratios have the same effect on the RT.

A uniform bill of materials composed of industry standard elements that are frequently reapplied throughout the construction footprint enforces cost control and reduces construction schedule delay risks. The selection of one standard cable sheath fiber count for all distribution branch runs and another for feeder runs between the LCP and RT is beneficial. The variables involved in selecting a distribution branch fiber count are NAP split ratio(s), anticipated home density range, HFC distribution branch length (typically on a per active basis) and an allocation for spare fiber. Table #1 illustrated how NAP split ratios can be selected based upon home density and branch length. In this case higher NAP port counts are implemented using two smaller optical splitters. Such an approach allows the deferral of the optical splitter and its associated cost.

Distribution Branch Cable Sheath Fiber Count Scaling			
Length	900' Typical		
Homes Passed	8 HP	16 HP	32 HP
CATV Tap Size	2 Port	4 Port	8 Port
Taps / Branch	4	4	4
NAP Split	1:1	2x1:2	2x1:4
PON Fibers	8	8	8
Spare Fibers	4	4	4
Total Fibers	12	12	12

Table #1

A generous allocation of spare fiber throughout the design will insure restoration and application flexibility in the future, and possibly aid in addressing legislative issues associated with equal-access. Adding individual fibers to sheath counts has a relatively minor impact on first cost. Later these fibers can be used to deliver high value multi-gigabit services via dedicated port connections from facilities remote to the subscriber such as a primary headend or hub.

At the time of this paper's authorship the industry does not provide a selection of low cost standard NAP enclosures suitable for co-location with standard and power-passing cable television taps. This condition will change as PON deployments become more routine.

Operations Engineering

Easy PON subscriber provisioning and low or no PON equipment maintenance are key operations cost control objectives. To meet these objective measures must be taken at each physical level of the PON as well as at the MAC protocol layer.

Installation of a subscriber service drop and CPE must be executable by a technician with a relatively low craft skill level and without the need of advanced tooling or diagnostic equipment.

Turn-up of the CPE must be a plug-n-play event with the MAC protocol automatically handling provisioning and network access issues in much the same way as done today by cable modems via the **data-over-cable service interface specification** (DOCSIS) protocol.

Optical budgets, noise and return loss margins must allow the use of industry standard mechanical fiber splice devices at both ends of the service drop. Experience has shown [3] that cable drop damage due to bending or dirty connections can occur during the installation process. The ability to easily open connections for **optical time-domain reflectometer** (OTDR) measurements at the NAP and below the last optical splitter is convenient. Additionally the PHY receivers at both at both ends of the PON must provide a dynamic input range that can tolerate the initial exclusion of the NAP's optical splitter.

The ability of the PON to tolerate the exclusion of the NAP splitter for the first NAP subscriber is more difficult to achieve when high optical split ratios are applied at the NAP.

The service technician cannot be required to access or open adjacent NAP or LCP enclosures to establish which PON fiber the subscriber's NAP has been allocated nor to activate or light the PON fiber. Having to access other devices during the provisioning process can be logistically difficult and is certain to significantly drive labor costs up on a per-home served basis. Again this objective is directly influenced by the allocations of splitter ratios at the NAP and LCP as well as cable sheath fiber count.

The LCP must be affordably provisioned such that each distribution cable sheath has

at least one dedicated fiber for each NAP on the sheath.

This fact coupled with a desire to keep first costs down encourages higher branch cable sheath fiber counts and lower LCP split ratios. During plant construction the appropriate fiber color for each individual NAP is tagged with color marker on or in the NAP, or alternatively the appropriate fiber can be ring-cut from the sheath and stored ready for use.

PON OVERLAY SYNCHRONIZATION

The proposed methodology of synchronizing a PON overlay with an existing or new HFC plant is driven by the objective of having the PON's key functional elements closely located, both physically and logically, with the corresponding HFC devices. Practically, direct physical co-location is required at the CPE, tap and RT levels. Strict co-location of the LCP with a HFC fiber node may be desirable, but not a requirement. The accomplishment of this objective must of course align with as many of the previously mentioned engineering constraints as possible.

The PON design is driven by three key parameters, the PON split ratio, the ratio's distribution through the network and the sheath/enclosure fiber count limits. These three parameters narrow the field of alternatives such that final selections can be made based upon specific performance or functionality objectives and first cost reduction measures.

PON Split Ratio Analysis

All possible ratio distributions between the NAP, LCP and RT were explored for the

combined split ratios of 1:1, 1:2, 1:4, 1:8, 1:16 and 1:32. The first three ratios offer limited flexibility with regard to distribution, but were explored for comparative purposes. The first column of Table #2 list the possible ratio combinations for each overall split ratio.

Limiting the associated fiber counts to a maximum of 136 controls the LCP's position or depth in the network. The LCP fiber counts shown in Table #2 are inclusive of both inbound and outbound fibers. With regard to split ratio allocation the LCP's position in the network is dominated by the NAP. For example, a 1:4 NAP ratio favors a 256 HP LCP position.

Alternatives can also be explored such as creating a 4 port NAP utilizing two 2 Port splitters at the NAP, this will favor a 128 HP position at the expense of doubling the number of fibers in each cable branch. At first such an alternative may be dismissed as costly, but individual fibers are relatively inexpensive compared to optical splitters. Having two fibers available at each NAP further delays the purchase of NAP splitters until 3 or more NAP ports are in service, which may never happen. The measure also reduces the dynamic range requirements imposed on the CPE and OLT PHY receivers.

Strategies that postpone the use of a splitter at the NAP significantly impact the average subscriber provisioning costs, particularly during the early years of plant operation while subscription rates are low.

Table #2

Split Combo NAP-LCP-RT	LCP / Physical Node		Rremote Terminal		Normalized	
	Size	Fibers	Size	Fibers	First Cost	Ultimate Cost
1-32 PON Split Ratio						
1-32-1	128 HP	132 F	4096 HP	128 F	2.2	2.2
1-16-2	128 HP	136 F	2048 HP	128 F	2.2	2.2
1-8-4	64 HP	72 F	1024 HP	128 F	1.4	1.4
1-4-8	64 HP	80 F	512 HP	128 F	1.4	1.4
1-2-16	64 HP	96 F	256 HP	128 F	1.4	1.4
1-1-32	64 HP	128 F	128 HP	128 F	1.5	1.5
2-16-1	256 HP	136 F	4096 HP	128 F	2.0	2.2
2-8-2	128 HP	72 F	2048 HP	128 F	1.2	1.3
2-4-4	128 HP	80 F	1024 HP	128 F	1.2	1.4
2-2-8	128 HP	96 F	512 HP	128 F	1.1	1.3
2-1-16	128 HP	128 F	256 HP	128 F	1.2	1.4
4-8-1	256 HP	72 F	4096 HP	128 F	1.0	1.3
4-4-2	256 HP	80 F	2048 HP	128 F	1.1	1.4
4-2-4	256 HP	96 F	1024 HP	128 F	1.1	1.4
4-1-8	256 HP	128 F	512 HP	128 F	1.1	1.4
8-4-1	512 HP	80 F	4096 HP	128 F	1.0	1.3
8-2-2	512 HP	96 F	2048 HP	128 F	1.0	1.3
8-1-4	512 HP	128 F	1024 HP	128 F	1.0	1.4
1-16 PON Split Ratio						
1-16-1	128 HP	136 F	2048 HP	128 F	2.2	2.2
1-8-2	64 HP	72 F	1024 HP	128 F	1.4	1.4
1-4-4	64 HP	80 F	512 HP	128 F	1.4	1.4
1-2-8	64 HP	96 F	256 HP	128 F	1.4	1.4
1-1-16	64 HP	128 F	128 HP	128 F	1.5	1.5
2-8-1	128 HP	72 F	2048 HP	128 F	1.2	1.3
2-4-2	128 HP	80 F	1024 HP	128 F	1.2	1.4
2-2-4	128 HP	96 F	512 HP	128 F	1.2	1.3
2-1-8	128 HP	128 F	256 HP	128 F	1.2	1.4
4-4-1	256 HP	80 F	2048 HP	128 F	1.1	1.4
4-2-2	256 HP	96 F	1024 HP	128 F	1.1	1.4
4-1-4	256 HP	128 F	512 HP	128 F	1.1	1.4
8-2-1	512 HP	96 F	2048 HP	128 F	1.0	1.3
8-1-2	512 HP	128 F	1024 HP	128 F	1.0	1.4
1-8 PON Split Ratio						
1-8-1	64 HP	72 F	1024 HP	128 F	1.4	1.4
1-4-2	64 HP	80 F	512 HP	128 F	1.4	1.4
1-2-4	64 HP	96 F	256 HP	128 F	1.4	1.4
1-1-8	64 HP	128 F	128 HP	128 F	1.6	1.6
2-4-1	128 HP	80 F	1024 HP	128 F	1.2	1.4
2-2-2	128 HP	96 F	512 HP	128 F	1.2	1.3
2-1-4	128 HP	128 F	256 HP	128 F	1.3	1.4
4-2-1	256 HP	96 F	1024 HP	128 F	1.1	1.4
4-1-2	256 HP	128 F	512 HP	128 F	1.1	1.4
8-1-1	512 HP	128 F	1024 HP	128 F	1.1	1.4
1-4 PON Split Ratio						
1-4-1	64 HP	80 F	512 HP	128 F	1.4	1.4
1-2-2	64 HP	96 F	256 HP	128 F	1.4	1.4
1-1-4	64 HP	128 F	128 HP	128 F	1.6	1.6
2-2-1	128 HP	96 F	512 HP	128 F	1.2	1.4
2-1-2	128 HP	128 F	256 HP	128 F	1.3	1.4
4-1-1	256 HP	128 F	512 HP	128 F	1.1	1.4
1-2 PON Split Ratio						
1-2-1	64 HP	96 F	256 HP	128 F	1.5	1.5
1-1-2	64 HP	128 F	128 HP	128 F	1.7	1.7
2-1-1	128 HP	128 F	256 HP	128 F	1.4	1.5
1-1 PON Split Ratio						
1-1-1	64 HP	128 F	128 HP	128 F	1.8	1.8

The RT is positioned in the network by following a process similar to that used with the LCP. The total incoming fiber count is

held constant by adjusting the number of LCP's reporting to the RT. As shown in the 5th column of Table #2 the total inbound LCP fiber count was held at 128. This quantity was chosen, rather than the previously mentioned 150, to allow for out-bound or ring fibers and HFC node fibers that were not included in the analysis.

Both NAP and LCP ratios influence the depth or position of the RT in the network. NAP/LCP ratios of 1:4/1:2 favor a 1000 HP RT in the 8, 16 and 32 PON splits.

Physical placement of ground mounted 1000 and 2000 HP RT cabinets is common, particularly with regard to digital loop carrier telephone systems. Small 500 HP RT's can be pole mounted, but become numerous and pose powering challenges when extended backup times are required to ensure network operation during extreme weather conditions.

The cost data provided in Table#2 is based on items that generate a relative cost difference between PON split options. For example, the fiber cost for a PON split of 1 – 1 – 32 would be considerably greater than 32 – 1 – 1 this along with other associated impacts such as number of RT's required for each combination. The resulting dollars per home passed were then normalized to the lowest first cost option. Normalization of the data is appropriate since the cost analysis is gauging relative merits rather than absolute values.

The cost impacts of various subscription take-rates are difficult to predict and are not included in the analysis. This is principally due to the cost of an optical splitter at the NAP if it is needed at the time of service provisioning. In many instances a splitter may not be needed for the early subscriber(s) whose drops can be connected directly to the PON fiber allocated to the NAP. Thus unless

there is a great deal of subscriber clustering the impact of NAP splitter costs will be negligible until take rates approach 30% to 50%. Due to this variability the analysis focuses on the initial relative cost and that of a 100% take-rate. The 100% figure reflects the comparative cost between low and high subscription rate pockets.

The principal conclusion drawn from the cost data is that from a relative-economics point of view a plant can be designed by appropriately applying a wide range of split combinations.

This is good news from a plant-engineering point of view, providing the network designer latitude to meet the numerous other non-cost related engineering constraints such as those previously outlined.

OBSERVATIONS

The key take-away observations from this analysis effort and discussion are:

1. Properly executed, PON optical split ratios do not greatly affect network first costs.
2. The appropriate allocation of PON split ratios between key network elements can be effectively used to control the physical and logical location of those elements in the network.
3. Control of sheath fiber counts can be effectively used to defer optical splitter cost at the NAP and to ease performance demands on PHY receivers and transmitters.
4. One or two standard sheath fiber count configurations can be readily adapted to varying subscriber densities and HFC tap configurations.
5. Mechanical fiber splice connections at the NAP and CPE are required to keep provisioning costs down and support drop

troubleshooting. The network's optical performance requirements must be tailored accordingly.

6. The eventual migration of the network to gigabit speeds is likely to occur earlier with lower overall PON split ratios.
7. The planned under subscription of a PON split ratio can be leveraged to reduce the cost of the RT while enabling lit fiber to pass by 100% of the serving area's subscriber base.
8. The utilization of a self-provisioning MAC protocol capable of dynamic bandwidth allocation without regard to PON split or subscription is required to maintain high customer satisfaction and low provisioning/maintenance cost.

This analysis was focused on PON split ratios in exponential powers of 2. Clearly other ratios are equally valid and have been used successfully [3] such as multiples of 3. A significant cost benefit would not be anticipated in alternative ratios, however they may offer useful alternatives with regard to LCP and RT placement and provisioning.

SUMMARY

Fundamentally, HFC alone is a proven financial workhorse capable of delivering outstanding residential video, voice and data services. Based upon the work done thus far it is clear that it will be some time before PON alone can compete economically in a standard residential market space with HFC. However, in part this fact is contrasted by three emerging near term trends:

- The growing utilization of e-business practices by small under-served businesses, businesses within easy service provisioning range of an HFC plant.
- The steady demand for advanced data networking and voice services for work-at-home and small branch office employees.

- A strong financial desire to reduce or eliminate the risk of obsolescence in newly constructed plant. Extending the plant's operation during its positive cash flow life by as little as two or three years has significant investment repercussions.

Perhaps the optimal migration compromise between the realities of today's economics and the future's service demands may be the strategic application of optical network overlaying a new or existing HFC plant.

Accordingly this paper has proposed a series of engineering considerations and design methodology for the synchronization of a PON with a modern HFC network. The methodology has been demonstrated analytically to be capable of effectively placing key PON elements at physical and logical locations consistent with an underlying cable TV network.

REFERENCES

- [1] P. Connolly, "New Carrier Class Architecture Busts Bandwidth Blues", NCTA Technical Papers, 2001
- [2] M. Conner, E. Dunton, "Future Proofing Your Residential Broadband Access Network", Technical Proceedings, 2000 National Fiber Optic Engineering Conference (NFOEC)
- [3] G. Mahony, "Fiber to the Home: A Summary of Bellsouth's First Office Application", Technical Proceedings, 2000 National Fiber Optic Engineering Conference (NFOEC)
- [4] "Broadband Optical Access Systems Based on Passive Optical Networks (PON)", ITU-T Recommendation G.983.1

BIOGRAPHY

Donald Sorenson is Manager of Advanced Access Planning in the Transmission Network Systems sector of Scientific-Atlanta, Inc.. Mr. Sorenson is responsible for the engineering evaluation and planning of Scientific-Atlanta's next generation access network technologies with particular emphasis on technology migration strategies for HFC networks. Donald obtained his BSEE from South Dakota School of Mines and Technology in 1978 and has maintained a professional engineering focus throughout his career. He brings 23 years of product and technical business development experience to S-A and the telecommunications industry, of which the past ten have been spent in the Cable Television and Telecom market segments.

email: donald.sorenson@sciatl.com

THE EXACT BER PERFORMANCE OF 256-QAM WITH RF CARRIER PHASE NOISE

Robert L. Howald, Ph.D.
Director, Transmission Networks Systems Engineering
Motorola Broadband Communications Sector

Abstract

The impact of phase noise on digital communications systems has been an ongoing challenge for system designers as modulations have become more complex, error corrections schemes have become highly sophisticated, and the channels under consideration become more varied. The HFC channel as used by CATV providers is a unique one from the phase jitter perspective, for a couple of key reasons. First, CATV is one of very few commercial uses of very higher order QAM, such as 64-QAM and 256-QAM, primarily because of the inherently high SNR and linearity due to the needed to support video. As a corollary to this point, as one of the few implementers of QAM, CATV is perhaps the only common example of a system that does so at RF frequencies. Secondly, precisely because the CATV plant was made with analog video in mind, the traditional channel was not designed with QAM in mind, including the equipment made for RF frequency generation. The result is upconverters and downconverters adequate for one application, and being asked to support another. While there is much to be thankful with respect to how digital signaling provides advantages all the way around, there are some possible "gotchas," and some overall confusion with respect to the phase noise topic. This paper is meant to describe and clarify these issues.

RAISING THE BAR

The use of 256-QAM has become increasingly popular in systems deploying digital signals in the forward band. There are several reasons for this. Most importantly, the expansion of the symbol set provides a 33% increase in bandwidth efficiency over 64-QAM. Also significantly, the SNR in the forward path is capable of supporting this sophisticated modulation format. Finally, the modem technology has become robust enough to handle the complexity of tracking, equalizing, and detecting 256-QAM symbols. As the constellation size has expanded, the sensitivity of the signal to forward path impairments has increased over and above the 6 dB difference in thermal noise sensitivity that exists between 64-QAM and 256-QAM. One of the impairments that can be particularly troublesome if not understood is RF carrier phase noise. This impairment is imposed when baseband digital data is modulated onto an RF carrier, and again during upconversion. On the receive side, it occurs again in the tuner in the settop box in the home, and in any subsequent stage that takes the tuner's IF output to baseband for processing by the receiver. Additional RF processing in between that involves frequency translation, such as block conversion, would also cause degradation.

There has been continual concern from the operator community about just how clean the RF carriers need to be to support upgrading digital services to 256-QAM.

Questions have arisen in some systems over whether problems encountered during its roll-out are related to phase noise issues. Phase noise has a mysterious aura about it that can cause confusion about what exactly it is, how it qualitatively impacts the transmission, and, most of all, how this translates to quantified performance degradation. For high levels of QAM in particular, the effects of phase noise have sometimes been crudely approximated, and sometimes been portrayed as intractable except via measurement. However, the impairment is in fact, completely analytically tractable, and degradation due to it can be precisely predicted. In this paper, we will present the approach to this analysis and develop the analytical solution. The results will be used to generate exact predictions in the form of 64-QAM and 256-QAM BER versus SNR versus phase noise curves. The analysis will be supported with plots and constellations to clarify the phase noise problem.

PHASE NOISE

Carrier Noise Spectrum

Digital modulations that encode information in carrier phase are naturally sensitive to impairments that disturb this phase. Carrier phase jitter, which occurs because of the inability in a practical implementation to generate ideal sinusoidal signals, is such a process. Figure 1 shows the nature of phase noise in the frequency domain, in this case using a Motorola C8U upconverter. An ideal sinusoid is represented by a line spectrum. In the figure, such an ideal sinusoid would be represented by a vertical line at 0 Hz on the x-axis, which is not defined for this logarithmic scale. Instead, carrier noise energy at frequencies away from the ideal line are displayed along the y-axis, with the offset frequency from zero on the x-axis. While just one "side" of the carrier noise spectrum is displayed, a phase

noise sideband at the same offset on the other side of the carrier has the same magnitude, and is anti-phase. There is a lower practical limit of offset frequency of interest in digital communications, because very slow phase noise as represented by very slow frequency variations are tracked at the receiver. In fact, as the zero frequency point is approached even more closely, the concept of phase noise is transformed into the concept of frequency stability, which characterizes the very slow frequency drift from the nominal center frequency associated with component and temperature variations of the circuitry. The amount that such carrier jitter disturbs detection is related to the QAM constellation's sensitivity to phase jitter, and the amount of this jitter that is imposed on the signal exists at the symbol detector.

As described above, there are two main components to the phase noise problem. Introduction of phase noise occurs during RF conversion. This includes both upconversion and downconversion. Figure 2 shows phase noise performance of a settop tuner, in this case tuned to 855 MHz. Note that the tuner's performance, as a low-cost, size-constrained, design, is significantly worse than the upconverter (Headend) piece. It is common for tuner performance to dominate the contributed RF phase noise. The phase noise spectrum at the tuner output represents essentially what is presented to the detection mechanism. The remaining IF conversion after tuning has minimal impact because of the effect of the upconverter and tuner, and because the downconversion from 44 MHz is a low frequency process, and thus less phase noisy.

The RF mixing that represents the points at which phase noise is introduced is one part of the two issues mentioned. The other part of this issue is the ability of the receiver to track the carrier phase. The amount of phase noise that is tracked is a spectral

function associated with the PLL filter processes involved in phase-locked carrier tracking. Basically, a PLL tracking loop follows the input reference noise roughly over its tracking loop bandwidth. Thus, it is incumbent upon receiver designers to understand the spectral mask of phase noise on the incoming signal, as well as the SNR characteristics of the received signal. A very wide bandwidth has advantages in its ability to acquire and track carrier phase relative to a narrowband tracking loop. The narrowband loop is less susceptible to false locking on spurious signals, and introduces less AWGN to the total phase error. The steady state tracking design is a trade-off of thermal noise degradation and phase noise degradation.

Quantifying Imposed Phase Noise

The quantification of the phase noise spectrum is determined by analyzing the frequency synthesizer design, which, in typical low cost commercial deployments, emphasizes low cost, direct divide PLL techniques. The performance of these PLL's can be modeled with computer aided mathematical analysis by understanding each of the noise sources and its contribution to the total noise performance.

Figure 3 represents an example of a simple synthesizer analysis that develops a phase noise prediction of a single-loop PLL synthesizer based on the relevant contributors to the output phase noise process. These include, primarily, the reference oscillator (crystal), the PLL logic devices, including phase detection and frequency division, and the voltage controlled oscillators (VCO's) that provide the RF output. Although this is a simple example, any level of synthesizer complexity it relatively easy to analyze, because PLL mathematics itself is quite straightforward.

On the receiver side, induced phase jitter is modified by the tracking mechanism. The receiver would include a tracking PLL, which could be analog or digitally implemented, depending at what level of receive downconversion is done. For single carrier modulations, times-N carrier recovery may be used, which still implements a PLL. In more advanced digital modems, decision-directed carrier recovery techniques may be used. The common Costas loop is a simple example. Regardless of mechanism, an untracked phase jitter component will exist at detection.

In Figure 4, the modification of carrier spectra and the resulting output error spectrum are shown. The output error spectrum in this case refers only to the difference between what phase noise spectrum is on the input, and what is on the output. While ideal phase noise would be to have none, an equally attractive case for performance would be one in which the receiver phase noise that is created by carrier tracking looks like that on the carrier, and thus the transmit and receive carrier references “move together,” nullifying error. When they are not together, we have untracked phase noise, and this is the impairment mechanism under discussion.

It is important to note that in Figure 4, both phase noise and thermal noise elements are shown to bring out the trade-off. Incoming phase noise is tracked up to the loop lowpass bandwidth, creating a highpass spectrum for untracked jitter. However, the additive thermal noise contributes to the loop phase noise as a lowpass function, limiting the close to carrier noise spectrum by creating this thermal noise floor. This then represents the design trade-off – low-passed thermal noise versus high-passed phase noise.

EFFECTS ON QAM

The other key element in understanding the phase noise phenomenon is to recognize the sensitivity of the underlying modulation scheme to the jitter. The need for bandwidth efficiency, ease of implementation, and high performance, makes the QAM family of modulations an excellent choice in many applications. The constellation is an important visual tool in evaluating modem performance for QAM modulations. Each distortion has its own characteristic effect on the constellation pattern. Consider Figure 5. For AWGN only, as in this 64-QAM example, the pattern will display a Gaussian "cloud" around the ideal constellation point. By contrast, phase noise represents angular rotation of the symbol points, away from the ideal as shown in Figure 6. The phase offset has a disproportionate effect on the outer symbol points, because the length the received point moves away from the ideal point of the constellation is proportional to the radius associated with that symbol. Consider a square, four-sided decision region that is $(2d \times 2d)$ in area, with an ideal constellation point in its center at coordinates (x, y) , and an error angle of ϕ due to phase noise like that shown in Figure 6. The shortest distance to each of the boundaries, without phase noise or any other impairment, is d in all directions. With the error ϕ , the four distances become modified, and shorter in two directions, to

$$\begin{aligned}d1 &= x \cos \phi - y \sin \phi - x + d \\d2 &= x + d - x \cos \phi + y \sin \phi \\d3 &= x \sin \phi + y \cos \phi - y + d \\d4 &= y + d - y \cos \phi - x \sin \phi.\end{aligned}$$

For M-QAM, the number of different symbol possibilities – the number of decision region squares – increases with M. Also, because of the relationship between the distance to a boundary and the symbol location on the constellation for a given

phase noise error, ϕ , higher density QAM signals become highly sensitive to phase noise. Outer symbol points are increasingly far from the center of the constellation, and this distance multiplies the phase error, shortening the distance to the boundary. Simple schemes such as BPSK and QPSK are quite robust to phase jitter. The symbol boundaries for QPSK are a full 45 degrees away. As such, a large error in phase would have to occur for a symbol error to be made in detection due to just jitter. Such a value of phase jitter, measured most conveniently in degrees rms, is extremely unlikely to occur with any significance to effect most reasonable error rates. This is not the case for QAM modulations using many more symbol points, and phase noise specifications must carefully consider the effects on the outer symbols.

How a noisy phase carrier effects symbol error rate performance is relatively straightforward to determine for the BPSK and QPSK situations under the assumption that the phase noise process is "slow". In this case, relative to a symbol time, the carrier phase can be considered constant over a symbol time, although its phase varies in value among the symbols. Small angle assumptions ($\sin \phi \approx \phi$) are often used to simplify BER performance analysis. Also, the phase noise is often considered to be from AWGN in the tracking PLL's only. In the case of HFC, both of these tend to be poor assumptions, because the larger M is, the more significant knowing the angle of error precisely becomes. For example, whereas the QPSK example pointed out that the symbols were 45° degrees away from a boundary, this number for 64-QAM is reduced to 7.7°, and for 256-QAM to 3.7° on the outer symbol points. In addition, considering just AWGN as the cause of jitter in the tracking PLL may work well for low SNR links, such as satellite or wireless. However, for HFC, the link SNR is quite high (40 dB range), so the need to account

for the other phase noise variables is significant. Assuming that AWGN-only jitter contributions exist leads to the conclusion that no error rate floor will be seen. This is a treacherous conclusion for multi-level M-QAM when true carrier phase noise is included.

Summarizing, then, the ability to overlook or be minimally concerned with phase noise impairments in many operating links today is a function of the robustness of the classical BPSK and QPSK families of modulations, and the low SNR channels they were often associated with. There is potential for a substantially larger penalty as the constellation is expanded.

PHASE NOISE STATISTICAL ASSUMPTIONS

The phase noise process is assumed to take on Gaussian behavior, and for practical implementation and mathematical tractability reasons, will be assumed to be both ergodic and stationary. Each of these latter two are common assumptions when developing phase noise analysis in the context of its performance effects for systems applications. This assumption is typically made in analysis of the phase jitter process when the goal is something other than exact calculation of the statistics of phase noise. In analysis for communication systems, the assumption of Gaussianity is assumed roughly equivalently as often as using the assumption that the PDF is Tikhonov. The Gaussian assumption is, in fact, a more general assumption than Tikhonov. The latter is derived from the assumption of a sine wave plus AWGN input only to the PLL (i.e. no phase noise). As mentioned, for a CATV situation and its high SNR, this is not a good assumption. In the limit of high SNR, the Tikhonov PDF asymptotically approaches a Gaussian PDF. Additionally, the Tikhonov PDF is arrived at by assuming a first order PLL. The basic

PLL structure here, and the dominant one in practice is the second order PLL topology.

The assumption of a Gaussian PDF leaves only the need to specify the two parameters that uniquely specify this distribution, the mean and the variance. A zero mean assumption (no steady state phase offset) leaves only the determination of variance, which represents phase jitter power. This can be determined by evaluating the phase noise spectrum over the region of interest that applies at symbol detection. This can be determined to some extent by phase noise measurements, although this type of equipment cannot easily decompose the spectrum into untracked and tracked phase jitter at a receiver. For design purposes, an accurate model is necessary. Such a model can be mathematically and empirically created by implementing the various PLL transfer and error functions, and applying expressions developed for the phase noise spectra of each of the various noise contributors into these models. Such an example was shown in Figures 3 and 4.

Untracked Phase Error

The spectrum of interest for systems analysis is that corresponding to untracked phase jitter, which is evaluated using the error function of the phase tracking process, the result of which is shown in Figure 4. Figure 4 has a highpass filter effect on the input phase noise mask. This highpass filter is the dashed frequency responses in Figure 4 at the top of the plot, and operates on an input spectrum such as Figure 2. Figure 4 looks similar to Figure 2 in the examples shown only coincidentally. The Figure 4 analysis assumed a thermal noise floor that creates the lowpass response, and the Figure 2 lowpass structure derives from high phase noise attributed to logic noise in the tuner design being amplified by the PLL. The tracking loop in the receiver will actually attenuate this region at 40 dB/decade if we

had visibility into that process, and the thermal noise added to the channel would instead result in a lowpass “floor” to replace it.

The spectrum derived from the model is used to calculate rms phase jitter by integrating the untracked jitter spectrum. The relationship between degrees rms and the composite integrated single sideband (SSB) phase noise in dBc (a negative dB value) is

$$\text{deg rms} = (180/\pi) \sqrt{10^{\phi(\text{dBc})}}.$$

The approach here, then, is to use the phase noise spectra developed above through mathematical and numerical modeling. The key parameters are the untracked rms jitter imposed at detection, and the bandwidth of interest over which it is imposed. The latter is a function of the signal bandwidth of interest at detection, which is basically tied to symbol rate.

BER EXPRESSION DEVELOPMENT

The BPSK and QPSK results have been derived in several classical reference papers. For constellations of higher order, less work has been done. However, these principles can be extended to larger constellations. For clarity, while this paper specifically addresses 256-QAM, the analysis will be explained using a much less cumbersome example, 16-QAM.

BER of 16-QAM with Phase Noise

The straightforward nature of QPSK analysis is due to the fact that each symbol is affected identically. Approaches to calculating degradation for larger constellations typically involved the small angle assumptions for sine and cosine. In this methodology, only the quadrature term contributes degradation and its variance can be added to the thermal Gaussian effects,

since they are uncorrelated. The end result is a composite SNR, calculated as

$$(1/\text{SNR}_T) = (1/\text{SNR}_{\text{awgn}}) + (1/\text{SNR}_{\phi}).$$

This approach is troublesome when signal amplitudes vary in M-QAM. The SNR approach breaks down because it does not properly account for the effect of angular rotation on the outer decision boundaries. It considers all symbols equally impacted by the jitter, effectively ignoring the sensitivity of the outer symbols, which limit the performance. For QPSK, all symbols have the same magnitude. For M-QAM, the crosstalk created by phase noise is a Gaussian noise term. However, it is multiplied by a discrete stochastic term (the data symbol) that can vary in magnitude, creating a different composite random variable. Fortunately, their likelihood is known, and thus a symbol-by-symbol decomposition is possible. Removal of the small angle simplification is therefore tractable for the slow phase noise case, so a complete solution can be derived.

For 16-QAM, there are three "classes" of symbols to consider. As such, it is most convenient to proceed with the derivation from the standpoint of the constellation, shown in Figure 7. There are the inner four symbols that form a QPSK constellation within the 16-QAM symbol set, the four symbols on the outer corners, and the remaining eight symbols which are on the outer ring. For two symbols a distance "d" apart (same d as previously used), the probability of selecting the wrong symbol when immersed in AWGN which is Gaussian of variance σ^2 is $Q(d/2\sigma)$, when optimal detection is used. Here, $Q(x)$ is the complementary error function

$$Q(x) = 1/\sqrt{2\pi} \int_x^\infty \exp(-y^2/2) dy.$$

This result is a simple conclusion resulting from determining the probability of a sample of Gaussian noise of standard deviation, σ , from exceeding $d/2$. Symbols bounded on multiple sides produce coefficient multipliers that increase this to $2Q(d/2\sigma)$, $3Q(d/2\sigma)$, or $4Q(d/2\sigma)$ in the AWGN case. It can also be shown that the $(d/2\sigma)$ argument is equal to $\sqrt{E_s/5N_o}$ for 16-QAM, which is also $\sqrt{4E_b/5N_o}$ for 16-QAM. Also, $SNR = E_s/N_o$ in a Nyquist channel.

Consider the first inner symbol, (a,a) , surrounded on all four sides, and dropping any $Q^2(x)$ terms that result from the analysis as negligibly small. For these inner symbols, attenuation and the crosstalk term are identical to that encountered in the QPSK situation, but the symbols in this case are bounded on both sides. Assuming the 16-QAM constellation shown previously, with constellation points at $(\pm a, \pm 3a)$, then consider the inner symbol (a, a) . This gives

$$P_i(\text{correct}) = \Pr[0 < a(\cos \phi + \sin \phi) + n_i < 2a] = P_q(\text{correct}).$$

Here, $P_i(\text{correct})$ is the probability of correct detection in the "I" direction (the x-axis), and $P_q(\text{correct})$ is the probability of correct detection in the "Q" direction, and n_i is an AWGN noise sample. This can be written

$$P_i(\text{correct}) = \Pr[-a(\cos \phi + \sin \phi) < n_i < 2a - a(\cos \phi + \sin \phi)].$$

Since the distance between points, $d = 2a$, this can be written

$$P_i(\text{correct}) = \Pr[-d/2(\cos \phi + \sin \phi) < n_i < d/2(2 - (\cos \phi + \sin \phi))].$$

Based on the previously stated results regarding the complementary error function, this can be expressed using $Q(x)$ and recognizing the following properties of the function:

$$\Pr(n_i > a) = Q(a/2\sigma)$$

$$\Pr(n_i < a) = 1 - Q(a/2\sigma)$$

and

$$Q(-x) = 1 - Q(x).$$

The result for is

$$P_i(\text{correct}) = (1 - Q[(d/2\sigma)(\cos \phi + \sin \phi)])(1 - Q[(d/2\sigma)(2 - (\cos \phi + \sin \phi))]).$$

Define

$$Q1 = Q[(d/2\sigma)(\cos \phi + \sin \phi)]$$

$$Q2 = Q[(d/2\sigma)(2 - (\cos \phi + \sin \phi))].$$

Again, neglecting the $Q^2(x)$ terms, negligibly small for practical error rates,

$$P_i(\text{correct}) = [1 - Q1 - Q2].$$

Similarly,

$$P_q(\text{correct}) = [1 - Q1 - Q2].$$

Then, the probability of a correct symbol, $\Pr(\text{correct})$, is

$$\Pr(\text{correct}) = P_i(\text{correct}) \cdot P_q(\text{correct}) = [1 - 2Q1 - 2Q2].$$

The error probability is then simply

$$P_e(a,a) = 2Q1 + 2Q2,$$

or

$$P_e(a,a) = 2Q[(d/2\sigma)(\cos \phi + \sin \phi)] + 2Q[(d/2\sigma)(2 - (\cos \phi + \sin \phi))].$$

This same error probability exists for the 16-QAM symbol $(-a,-a)$. These two symbols make up an eighth of the total symbol set,

which means that the contribution of these symbols to the overall 16-QAM degradation can be written

$$P_e = \frac{1}{4} Q[\sqrt{4E_b/5N_0}(\cos \phi + \sin \phi)] + \frac{1}{4} Q[\sqrt{4E_b/5N_0}(2 - (\cos \phi + \sin \phi))].$$

For the other two inner symbols (a,-a) and (-a,a), the only piece of the above analysis that changes is the sign in the crosstalk term. In the case here, that sign is reversed.

All of the symbols can be analyzed in this way, and various sets of symbols are impacted identically in terms of mathematical evaluation, with minor changes in signs and coefficients. It is also important to keep in mind that the result of averaging all of the symbol possibilities (each has 1/16 probability of being transmission) is a symbol error probability. A bit error probability is obtained by assuming the symbol error probability means the incorrectly selected symbol is a neighbor, and neighbors differ by one bit if properly (Gray) coded. This

assumption is not a good one for all cases of symbol error, particularly those that are due to catastrophic events such as large interference or sync loss.

For a complete BER expression, then, this conditional probability is then averaged over the Gaussian PDF of the untracked phase jitter and numerically evaluated.

RESULTS

Figure 8 shows the uncorrected BER results for 256-QAM with untracked phase noise varying from .5° rms to 1° rms. Clearly, the sensitivity of 256-QAM to this impairment is on display. A good rule to remember for rms jitter is that a signal-to-phase noise ratio of 35 dB represents 1° rms, and the degree rms term is a “voltage”, meaning .5° rms is a 41 dB signal-to-phase noise ratio. Clearly,

with this 41 dB ratio, there is trouble at hand. At a BER of 1E-8, there is a 2 dB degradation, which is generally unacceptable. Just as clearly, more than .5° rms quickly becomes intolerable.

Figure 9 shows the performance for values leading up to .5° rms. It is apparent from this plot that the BER curve is beginning to “take off” and flatten out at .5° rms, when compared to .35° rms, which shows a degradation less than one dB. An untracked phase jitter specification in the .25-.35° rms shows a tolerable situation.

Figures 10 and 11 graphically display why 256-QAM becomes so touchy to increased phase jitter. The effect of the additional jitter is magnified on the outer symbol points, creating a more troublesome scenario for the .5° rms case. The outer symbols for this case extend significantly closer to the boundaries than do the .25° rms case, making them more likely to have an AWGN sample push them over the decision boundary.

Figure 12 shows an example of measured case that represents about .5° rms of untracked jitter. Test and verification of phase noise effects is not an easy thing to do, because generating specific amounts of phase noise is not straightforward. While it is easy to create a noise local oscillator by adding noise to its control input, this is not the same mechanism by which phase noise is generated in real oscillators. Thus, the phase noise spectrum may look similar, but for different reasons. However, by statistically assuming Gaussian behavior, statistical differences are likely to be less pronounced than by focusing on deriving exact PDF's.

In this case, there was no noise injection for phase noise, just a cascade of RF sources as would be part of an HFC system – a Motorola C6U upconverter, and Temic tuner

in front of a Broadcomm test demodulator. The untracked jitter number is an estimate based on the imposed phase noise measured due to upconverter and downconverter (tuner) together, and knowledge and characterization of the test demodulator's tracking loop. Figure 2 shows a table of integrated phase noise values. This table breaks up the rms jitter contributions into spectral regions, and these regions combined with tracking bandwidth information lead to estimates of untracked jitter. As an example of the meaning of the rule of thumb previously given, the region between 1 kHz and 10 kHz from the carrier shows -31 dBc integrated noise power, or a signal-to-phase noise of 31 dB, which is about 1.6° rms.

With a 35 dB SNR in the 256-QAM channel (as measured by spectrum analyzer) the measured BER in the example of Figure 12 is about $3.7\text{E-}7$, against theory that predicts about $9\text{E-}8$. As a favorable sign, the realized BER measurement is slightly worse. In this case, the theory and measurement differ by about only .75 dB SNR. Note that the scale of E_b/N_0 and SNR are the same, merely offset by $10 \log 8 = 9$ dB, because 256-QAM contains eight bits per symbol.

For comparison a 64-QAM example using the same RF setup is shown. Figure 13 shows the expectation of 64-QAM as a function of phase noise, and points out that $.5^\circ$ rms is an acceptable value in this case, whereas 1° rms begins to become problematic. This example runs error free, because of both the high SNR and the (relatively) low phase noise imposed. Figure 14 shows through the broad clear area around the constellation points why this is so.

Error Correction

Sophisticated forward error correction (FEC), as is used on the forward path for digital video and data, has the powerful

ability to correct many forward path sins, and phase noise is among them. However, phase noise is a burst error mechanism because the phase error varies generally very slowly relative to the symbol rate. In other words, the jitter energy has frequency content much lower than the symbol rate. Thus, if the phase error drifts enough to cause a symbol error, it is likely that it will cause multiple symbol errors for a sequence of symbols. The amount likely is related to the jitter energy bandwidth relative to the symbol rate. Of course, it is not desirable to utilize the FEC capability to clean up phase noise issues, as this leaves less "margin" to handle other simultaneous impairments that the FEC was actually implemented to handle.

CONCLUSION

The use of 256-QAM places stringent phase noise requirements on any RF system. In the case of HFC, the issue is helped by the fact that the channel in question is very high SNR, leaving significant flexibility in the design of carrier tracking loops, such that imposed phase noise by converters can be removed at detection. The nature of the systems today is that the performance of converters and tuners can be as such to create bit errors, although not enough to cause catastrophic error problems that can be characteristic of higher order QAM in phase noisy channels. In general, FEC will correct for this, but at the expense of less FEC power available for other problems. More recently, lower phase noise equipment designs, made specifically to address high order QAM systems, have eased concerns about phase noise impairments. A complete analysis of a quality upconverter and a low phase noise tuner, coupled with a relatively wideband tracking receiver, can show that there adequate phase noise performance for 256-QAM transmission. The same analysis will also show that older equipment may yield acceptable, although impaired,

uncorrected BER results. However, the game once again changes if transmission standards evolve to 1024-QAM. A rough analysis shows that this modulation will be looking for better than $.1^\circ$ untracked rms jitter.

References

Howald, R., The Communications performance of Single-carrier and Multi-carrier Quadrature Amplitude Modulation In RF Carrier Phase Noise, UMI Dissertation Services, Ann Arbor, MI, 1998.

C8U-14



.25° rms > 1 kHz

Figure 1 – RF Carrier Phase Noise Imposed by C8U at 750 MHz

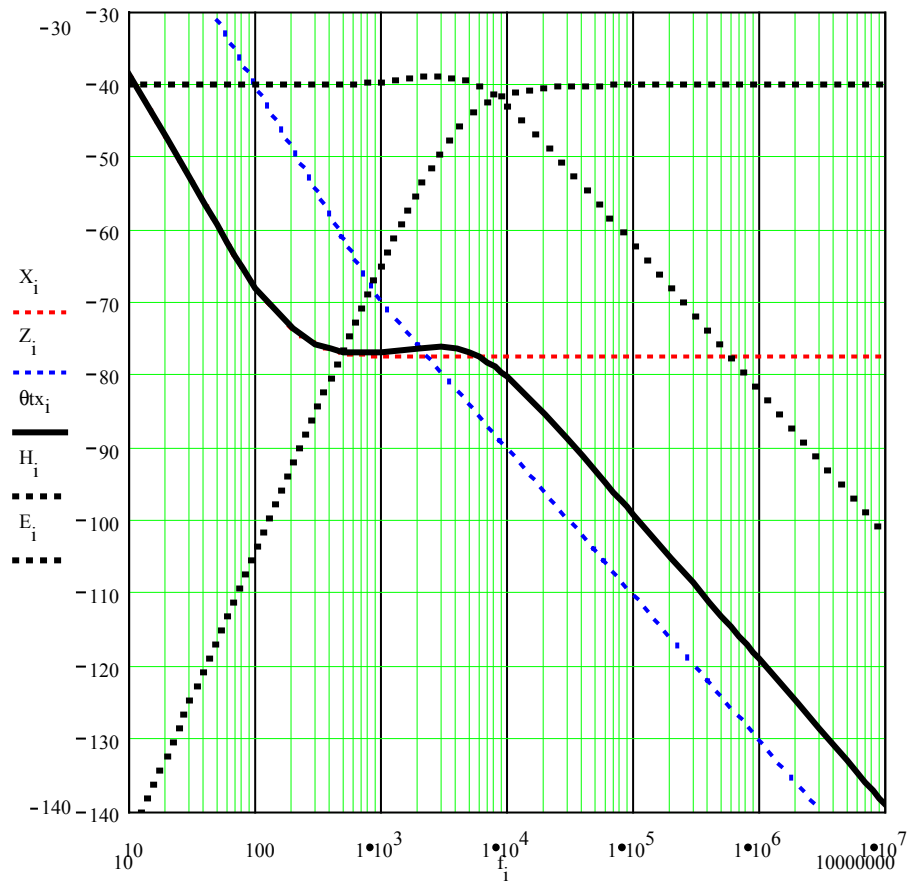


INTEGRATED NOISE LIST

START FREQ	STOP FREQ	INTEG NOISE	INTEG NOISE-SPURS
10.0 Hz	100.0 Hz	-10.0	-10.2
100.0 Hz	1000.0 Hz	-5.5	-8.5
1000.0 Hz	10000.0 Hz	-31.2	-31.4
10000.0 Hz	100000.0 Hz	-43.1	-43.3
100000.0 Hz	1000000.0 Hz	-48.0	-48.3

Figure 2 – RF Carrier Phase Noise Imposed by Tuner at 855 MHz

**Closed Loop PLL Response, Inside Loop BW Noise,
Oscillator Noise, and Total Phase Noise Response**



**Figure 3 – Example Development of RF Carrier Phase Noise
Spectrum from Synthesized Source**

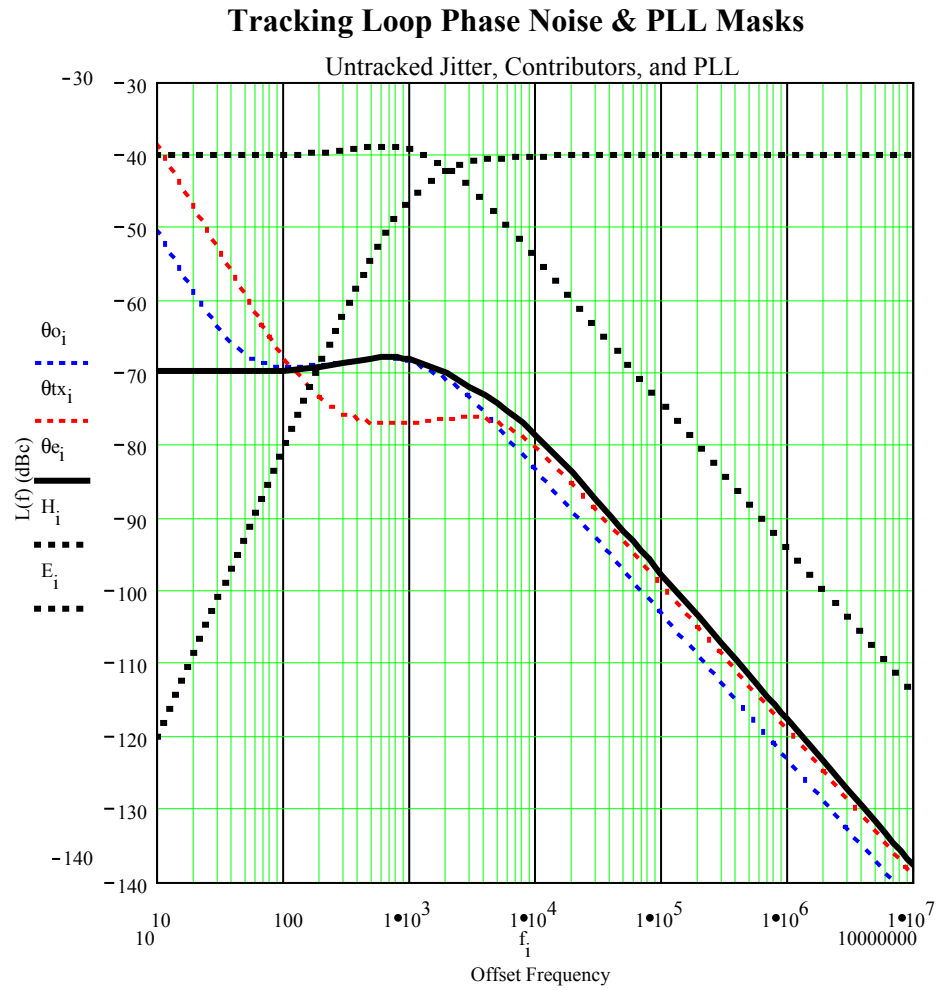


Figure 4 – Example Development of Demod Input Untracked Phase Noise Spectrum of AWGN & ϕ_n

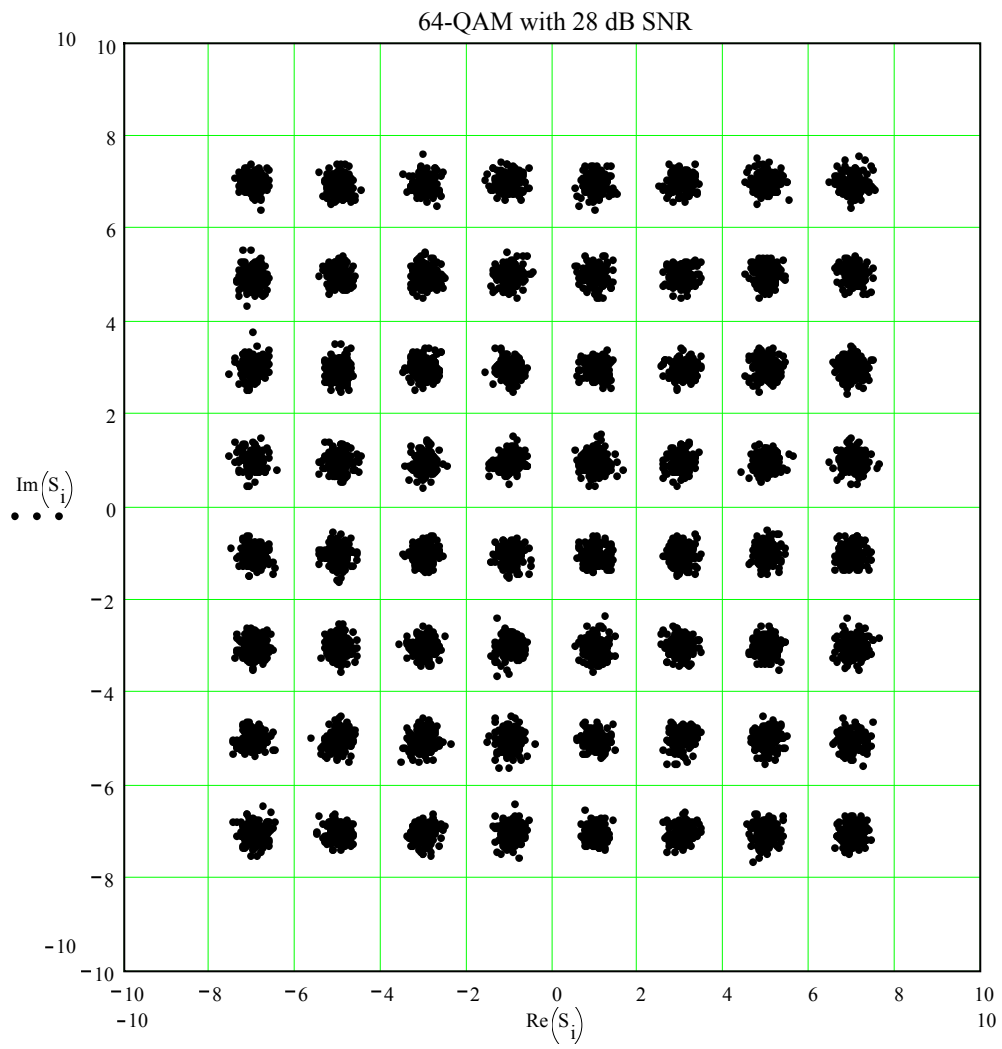


Figure 5 – 64-QAM with 28 dB SNR (BER ~ 1E-8)

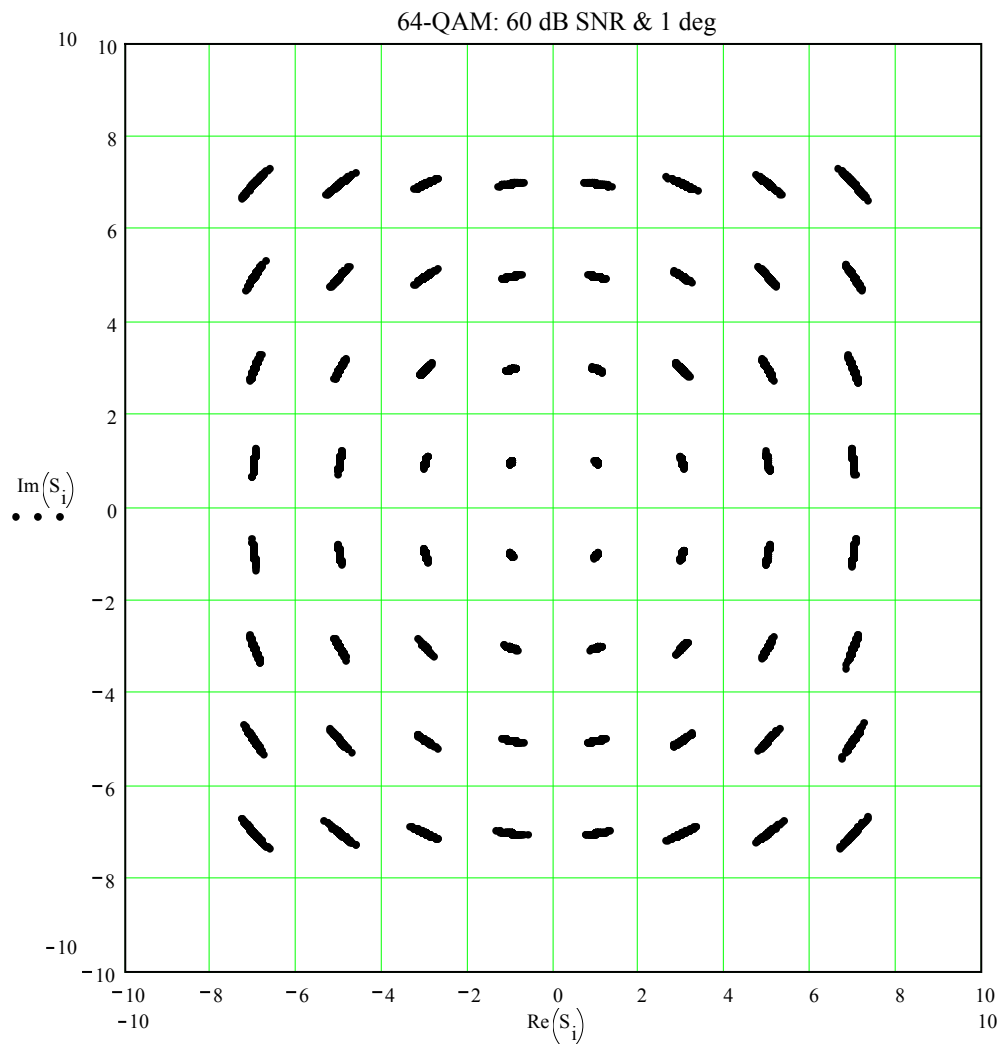


Figure 6 – 64-QAM with 60 dB SNR and 1 deg rms Phase Noise

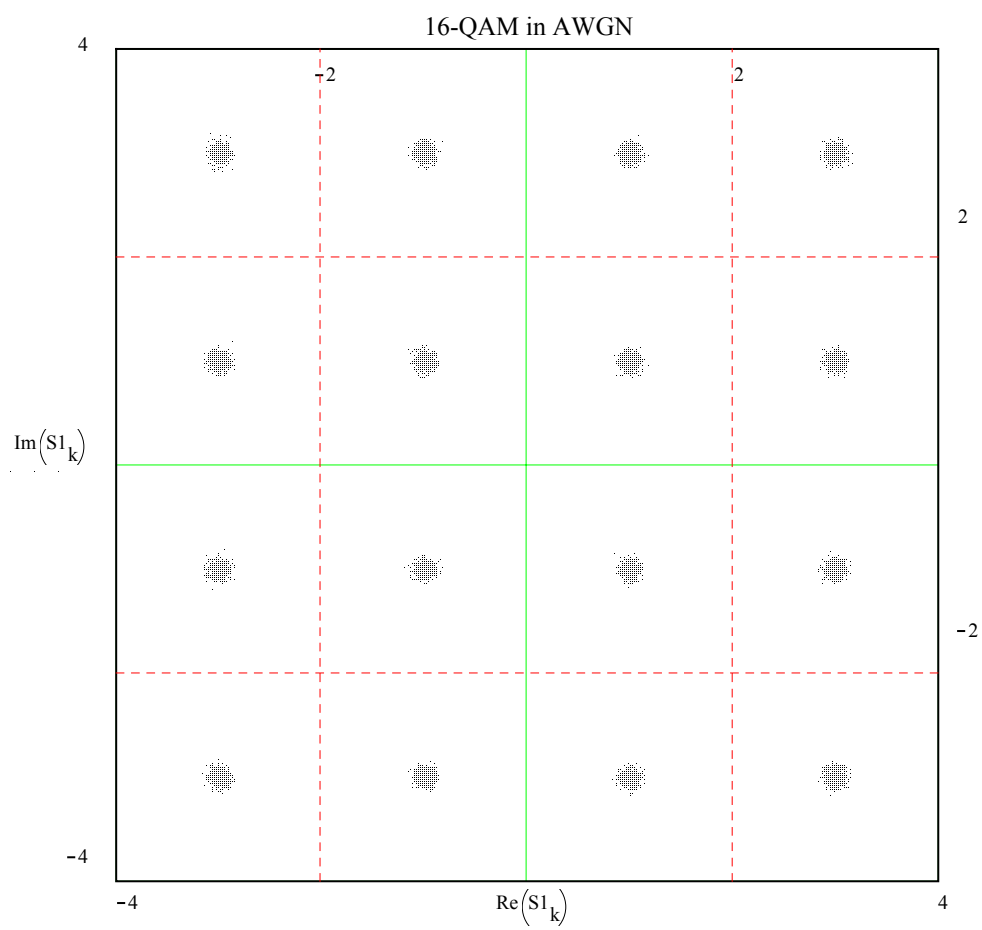


Figure 7 – 16-QAM in AWGN (SNR = 33 dB)

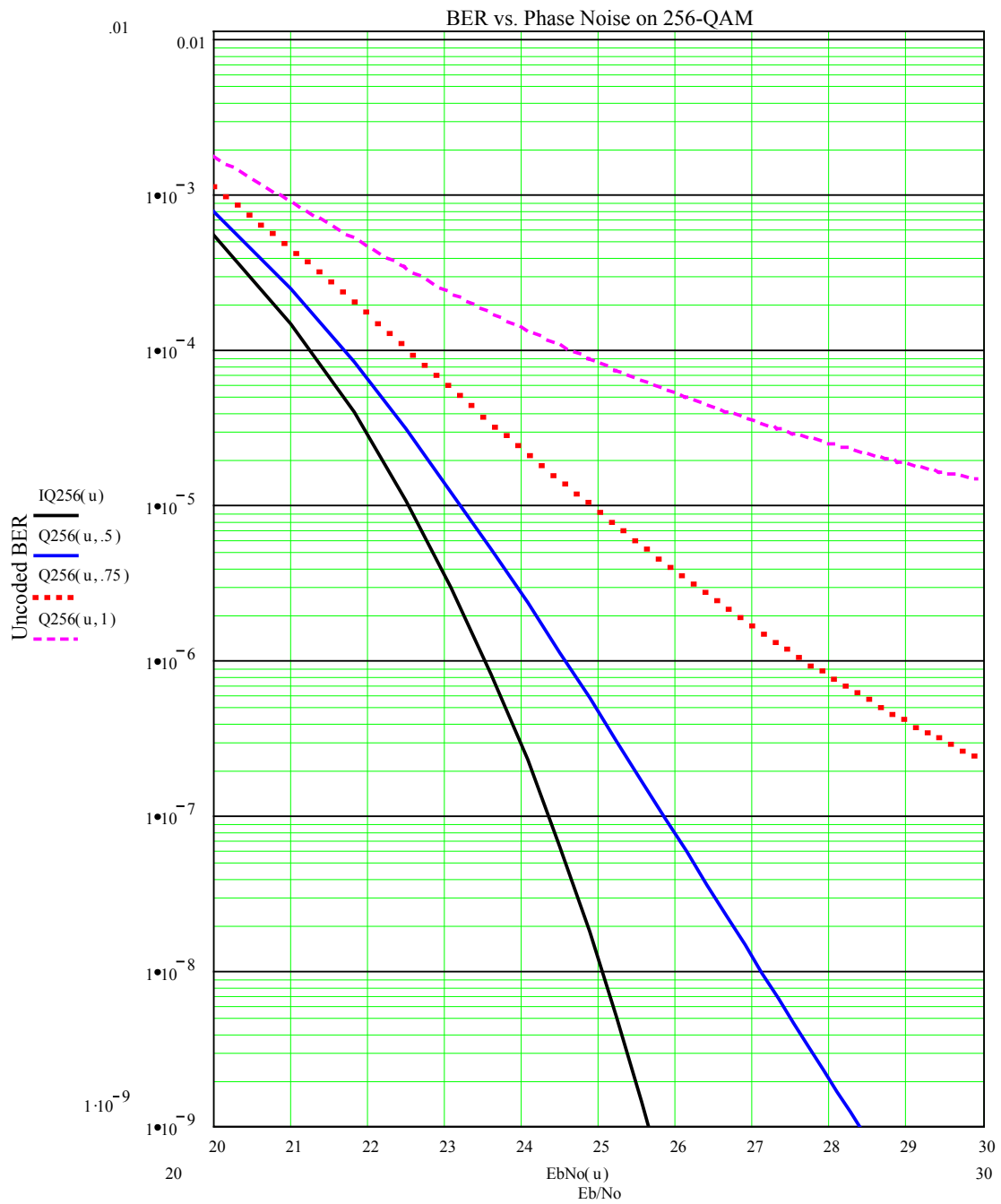


Figure 8 - 256-QAM BER with .5°, .75° and 1° rms of Gaussian Untracked Phase Jitter

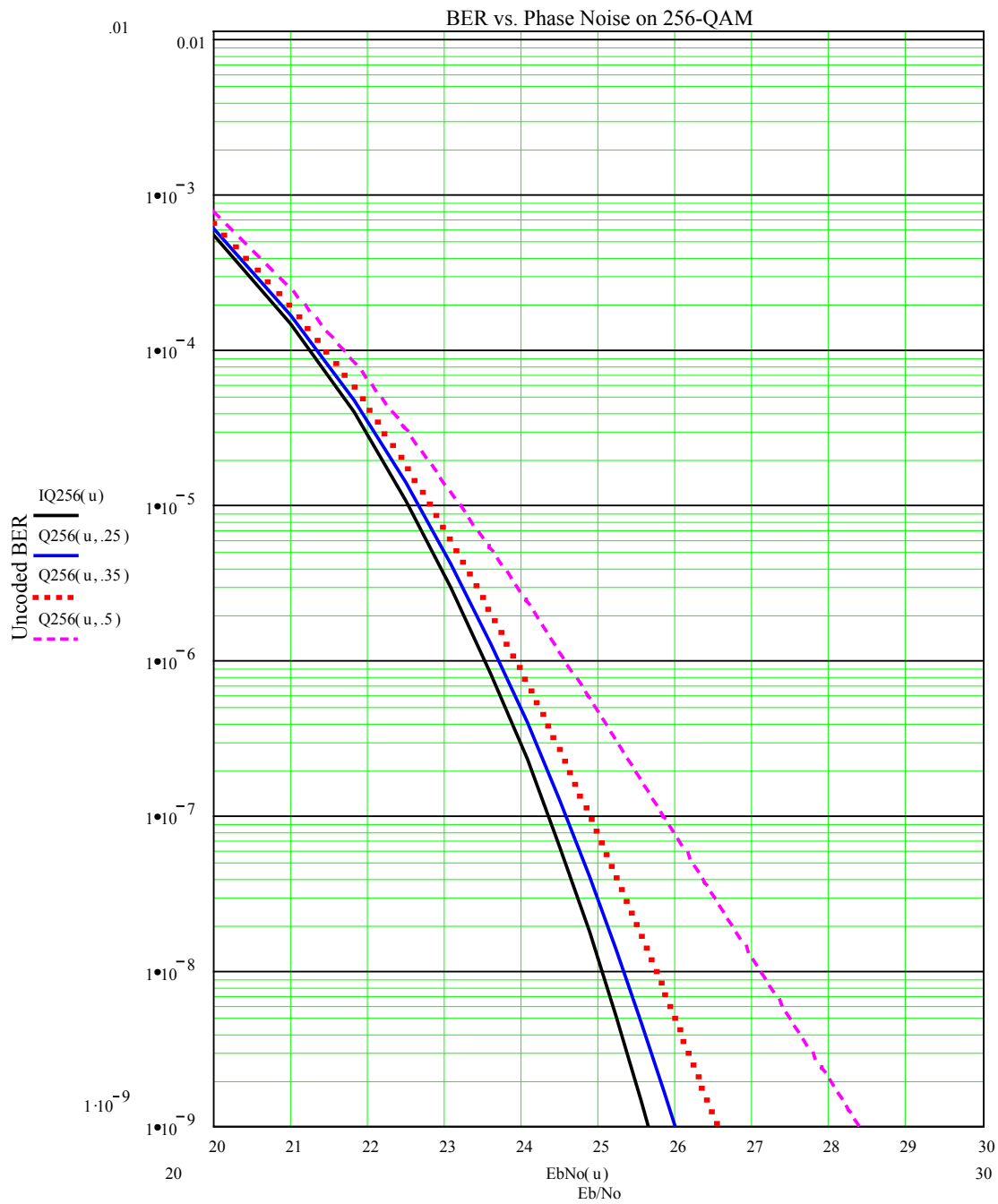


Figure 9 - 256-QAM BER with .25°, .35° and .5° rms of Gaussian Untracked Phase Jitter

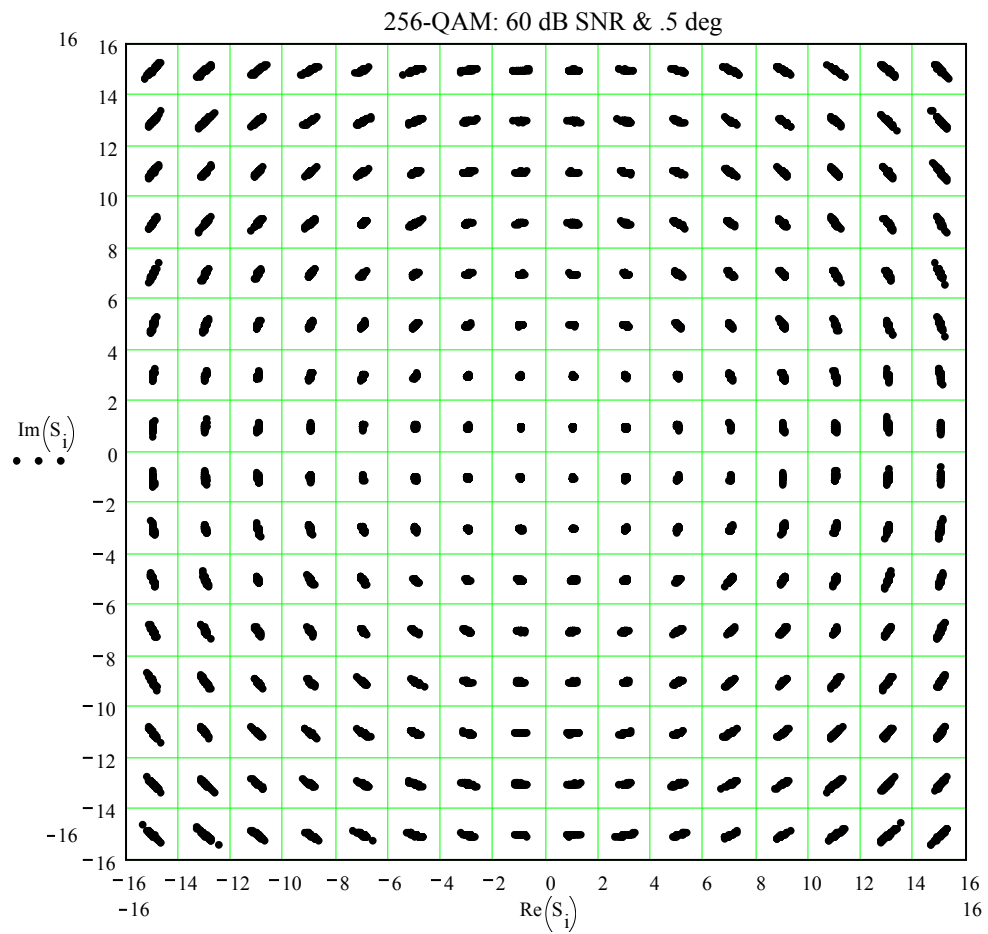


Figure 10 - 256-QAM Constellation with .5° rms Untracked Phase Jitter

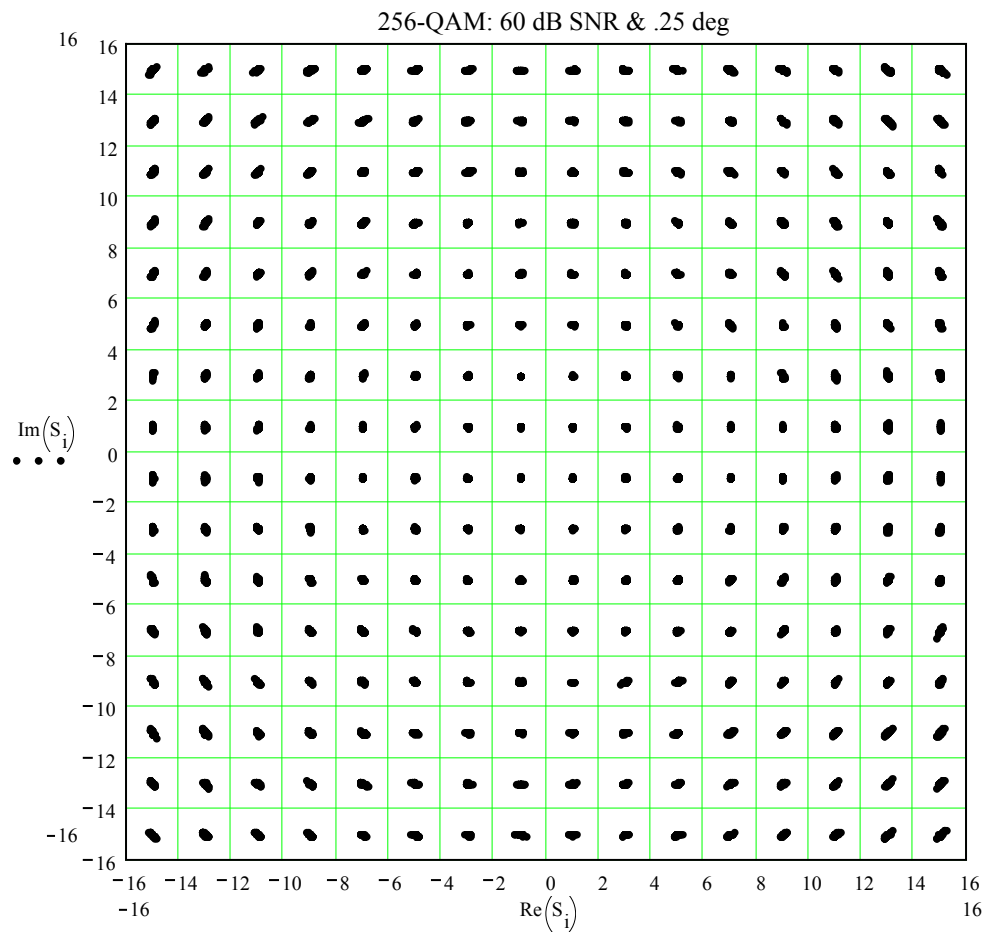
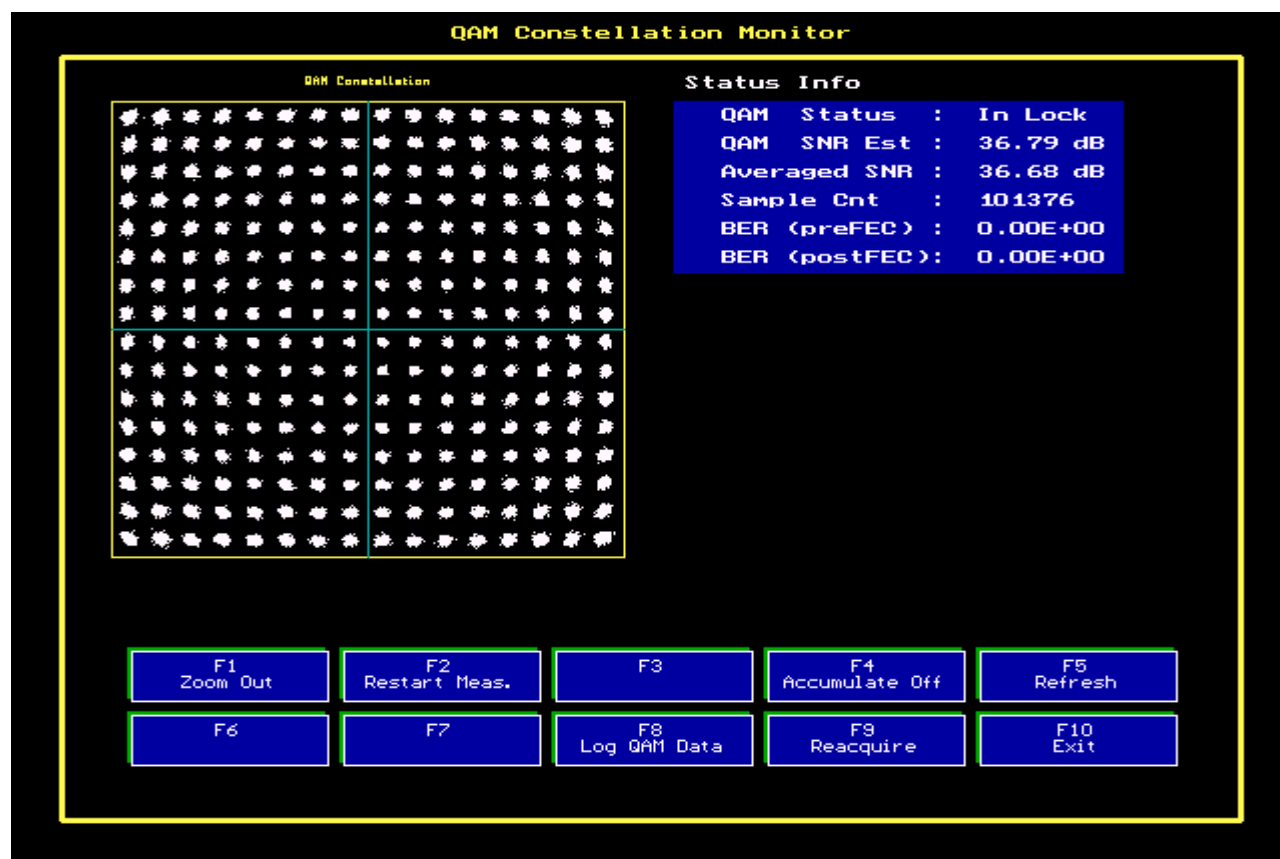


Figure 11 - 256-QAM Constellation with .25° rms Untracked Phase Jitter

256-QAM @ 35 dB SNR ($E_b/N_0 = 26$ dB), 600 MHz (approx .5 deg rms)



Predicted BER via Analysis = $9 \text{ E-}8$

Measured BER = $3.7 \text{ E-}7$

SNR Error Delta ~ .75 dB

Figure 12 – Measured 256-QAM (SNR = 35 dB) with .5° rms Untracked Phase Jitter

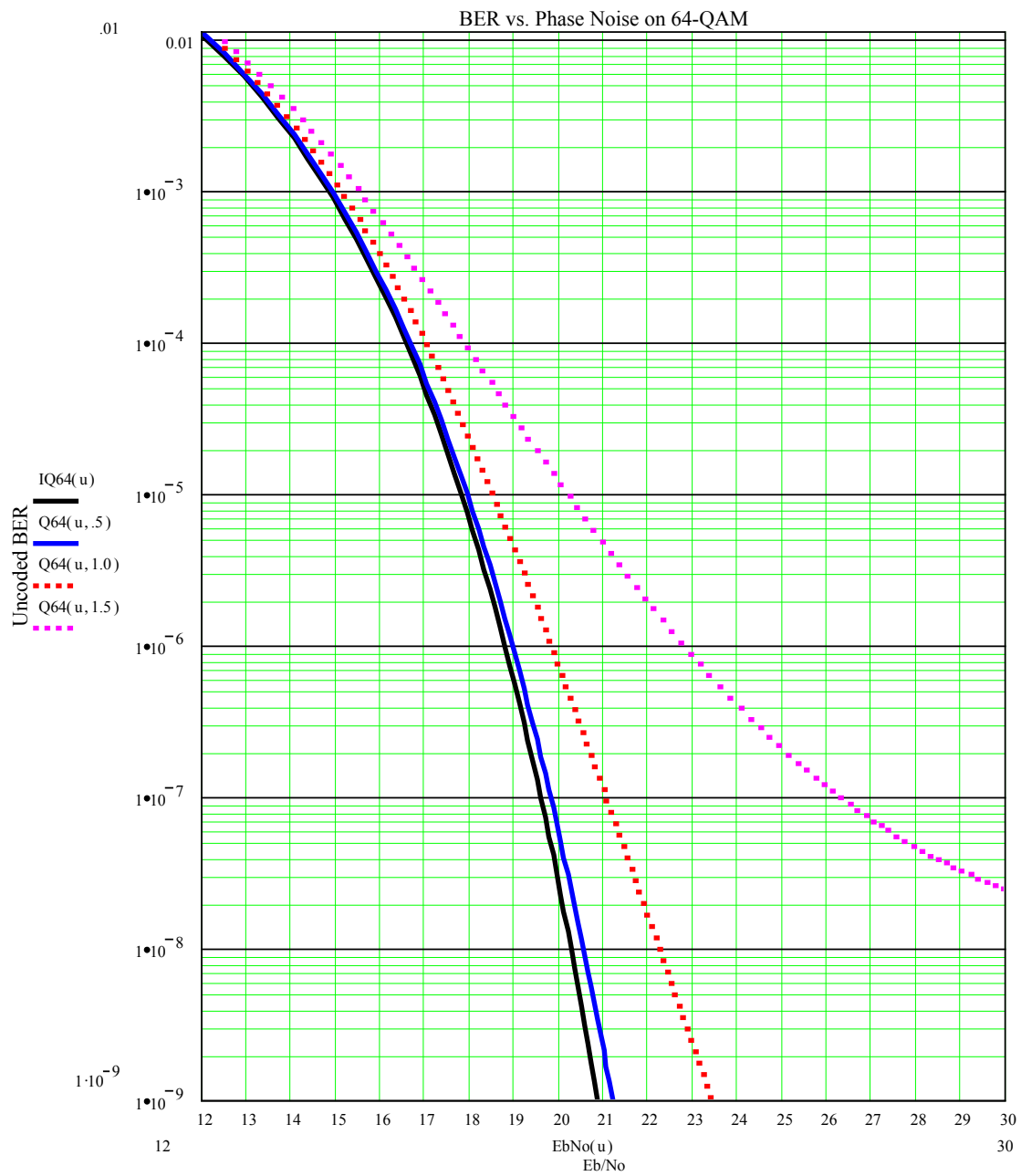


Figure 13 - 64-QAM BER with .5°, .1° and 1.5° rms of Gaussian Untracked Phase Jitter

64-QAM @ 35 dB SNR ($E_b/N_0 \approx 27$ dB), 600 MHz (approx .5 deg rms)

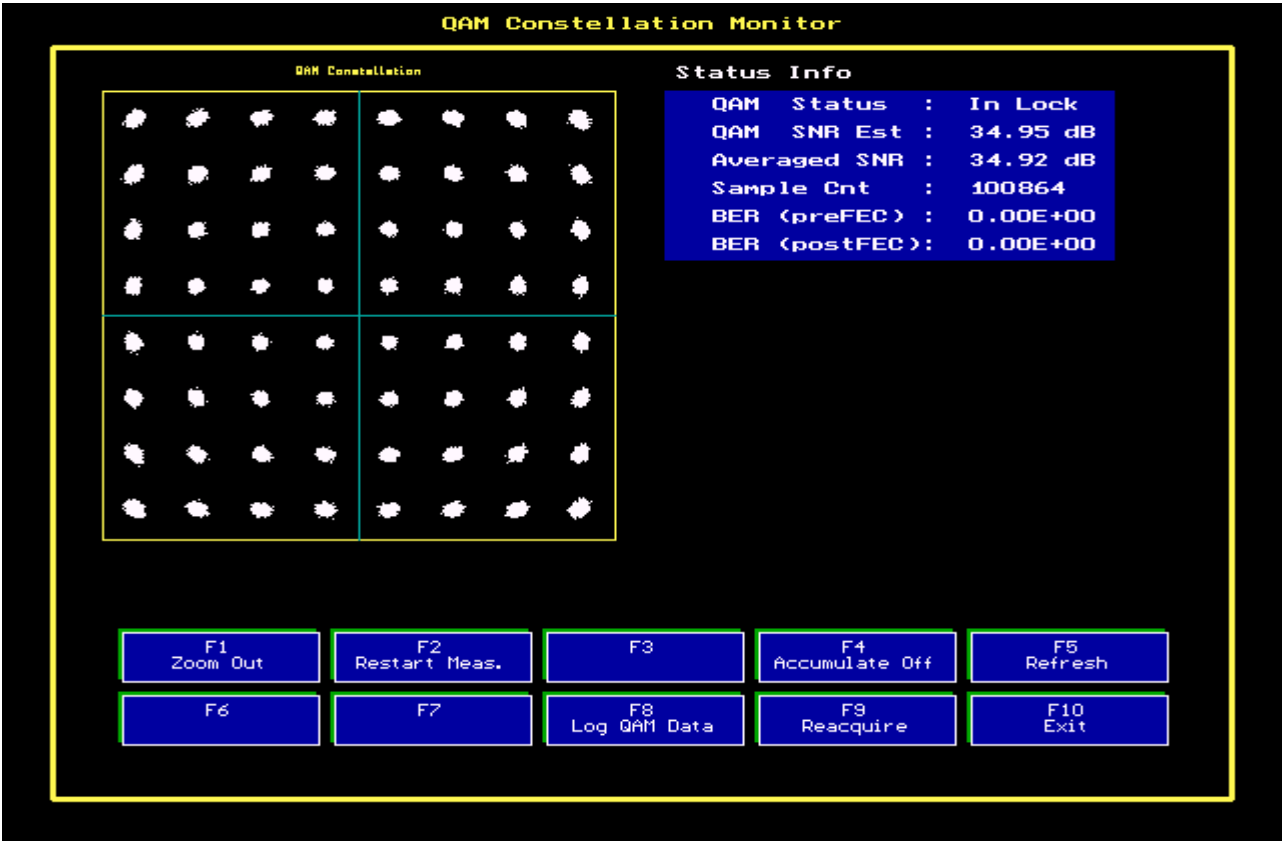


Figure 14 – Measured 64 -QAM (SNR = 35 dB) with .5° rms Untracked Phase Jitter

THE FUTURE OF STREAMING MEDIA

Marc L. Tayer
Aerocast, Inc.

Abstract

The final decade of the 20th Century was a very exciting one for the North American cable television industry. With major video threats emerging from the Direct To Home (DTH) sector, cable responded quite effectively with its digital cable rollout. Furthermore, with the Internet rapidly becoming a “must have” across America, cable took the high road with the successful standardization of DOCSIS along with market leadership in the broadband segment. More recent initiatives such as HFC telephony, VOD and HDTV are now beginning to gain momentum. Extensive fiber deployment and upgrades to two-way capability became the technical underpinnings for many of these new services.

Broadband streaming media is the next big wave and will create dramatic shifts in the way content is created, delivered to, and viewed by subscribers, affecting each constituency along the way. Moreover, broadband streaming media will enable an entirely new category of content, offering limitless possibilities and unprecedented choice for consumers. Operators who embrace broadband streaming will not only enhance their competitive market position, but will also find new revenue opportunities.

This paper provides an overview of the broadband streaming media technology and opportunity, contrasting it to the now familiar digital television experience, and including a brief discussion of the various emerging standards and technical elements which make up the broadband streaming media systems of the future.

OVERVIEW

The advent of digital television in the 1990s revolutionized the television industry by rapidly multiplying the choice of content available to consumers of DTH satellite and digital cable systems. In fact small dish DTH services such as DirecTV would not have been economically viable in the United States without digital television technology since the content offering from a single 32 transponder analog Ku-band Broadcast Satellite Service (BSS) satellite would have been insufficient to attract more than a niche subscriber base.

Nearly a decade later, with the maturation of digital television occurring due to successful standards such as MPEG-2, increasingly sophisticated and cost-effective set-top boxes, improved customer service and, most importantly, the creation of a critical mass of new programming services leveraging the economic bandwidth benefits of digital compression, the question becomes “what is next?” While MPEG-2 based Video On Demand (VOD), Personal Video Recorders (PVRs) and HDTV are emerging as compelling new offerings, they are evolutionary advances relative to what will become the single biggest leap in video communications since the invention of television itself: broadband streaming media.

Defined as entertainment-quality video delivered over Internet Protocol (IP) networks, enabling content to be viewed in near-real-time relative to the server’s stream, broadband streaming media will finally fulfill the consumer’s ultimate dream of viewing any content at any time. Broadening and further enhancing this category, a new

generation of IP PVRs (integrated into PCs and advanced set-tops) will enable downloadable media and the creation of home networked digital content libraries, making VHS collections, CDs and even DVDs seem as archaic as the LP shelves of the early 1980s.

A confluence of the following factors will bring forth an inexorable broadband streaming media content tsunami:

1. Rapidly Declining Bandwidth Costs
2. MPEG-4 and a New Generation of Open Standards
3. Secure Digital Rights Management
4. Intelligent Content Distribution Techniques
5. Proliferation of Broadband IP Modem Devices
6. Cost Effective and User Friendly Hardware and Software Tools

Rapidly Declining Bandwidth Costs

Today it is uneconomical to deliver broadband, and arguably even narrowband, video over the Internet. With Internet backbone bandwidth and content delivery costs in the range of \$500-2000 per Mbps/month (based on peak usage measurements), companies currently delivering IP video to consumers are doing so primarily to promote something else. In most cases, the video is really a form of advertisement or promotion in which the content provider is incurring a cost (or, more optimistically, making an investment) for another purpose, such as keeping eyeballs on a Website, promoting brand or enhancing product awareness, as opposed to providing the video as a product in and of itself.

For example, at \$2000/Mbps/month, assuming roughly 50% bandwidth utilization, a 2 hour movie coded at 750 kbps would cost nearly \$7 per viewer just for transport. This is too high by a factor of at least 10X.

Similarly, with the same transport pricing for an ad-supported 10 minute video encoded at 500 kbps (including roughly a minute of advertising at \$35 Cost Per Thousand viewers, or CPM), the pricing is at least 5X too high for economic viability.

Fortunately, within the next few years it is likely that Internet backbone bandwidth costs will decline to less than \$100 per Mbps/month. At these rates, the cost of delivery on a unicast basis will no longer be an economic impediment, whether the content is offered on an advertising, subscription or pay-per-view basis. Furthermore, Internet bandwidth alternatives such as charging on a per Megabyte delivered basis are becoming increasingly common.

The implicit bandwidth cost of the last mile operator must also be considered. In a digital world, every bit counts and it is logical and natural for a cable operator to allocate its bandwidth toward the maximum cash flow dollars. With VOD becoming an increasingly important product offering for operators, both to generate additional revenue and to differentiate their services from DTH operators (which have been limited architecturally to NVOD but are now wisely deploying PVRs to attempt to close the gap), streaming IP video will ultimately prevail as the predominant VOD technology. MPEG-4 over IP has such substantial bandwidth advantages over traditional MPEG-2 VOD, further strengthened by the immeasurable benefits of being part of the global IP network, that it will become the bread and butter of tomorrow's consumers' media appetite.

MPEG-4 and a New Generation of Open Standards

For broadcast digital television, one of the attractive features of the MPEG-2 video standard is its asymmetric architecture, with complex encoders and relatively simple (i.e., low cost) decoders. Since the MPEG-2 standard doesn't specify how to build an encoder (so long as the resulting bitstream syntax is MPEG-2 compliant), this design allowed MPEG-2 encoder companies to continually improve their algorithms (and therefore bandwidth efficiency for the content providers) in a fully compatible way with MPEG-2 set-top boxes. In addition, innovations such as statistical multiplexing, which continuously re-allocates bits between multiple video signals sharing a common bitstream pool, allows today's content providers to pack numerous video signals into a single satellite transponder or 6 MHz cable channel. Finally, the asymmetric design of hardware encoders did not impede market growth since the finite shelf space for content on DTH and cable television systems resulted in relatively few programmers leveraging a subscriber base in the tens of millions.

Yet, as far as MPEG-2 has come, it is now essentially a mature technology. For a CCIR-601 half resolution picture (352 x 480) of film material, coding rates of 1.5-2 Mbps are state-of-the-art to achieve entertainment quality, even using the best algorithmic tricks. For HDTV (1080i), good quality can be achieved at rates as low as 12-15 Mbps, depending on the content. And DVDs, using Variable Bit Rate (VBR) coded MPEG-2, typically employ average coding rates of 5-6 Mbps for CCIR-601 resolution.

MPEG-4 is now emerging as the standard for low bit rate video at rates less than 1 Mbps. Using the Advanced Simple Profile, a well-implemented MPEG-4 encoder will enable full screen entertainment quality video

at 750 kbps, perhaps even lower depending on the content. This startling level of bit rate efficiency will dramatically alter the VOD landscape. Since VOD implies an individual stream per user ("unicast" in IP parlance), the phrase "every bit counts" takes on an even more critical meaning versus a broadcast ("multicast") world where a few hundred kilobits per second of coding inefficiency is much less problematic since thousands or even millions of subscribers are tuning in to the same stream.

As with MPEG-2, various MPEG-4 profiles are being established to ensure interoperability. For example, in addition to the Advanced Simple Profile, there is the Simple Profile, the Core Profile, the Main Profile, and a few different Scaleable Profiles. The Advanced Simple Profile appears to be the best profile for streaming media applications, and the Scaleable Profiles will likely experience increasing use over time.

Beyond its impressive bit rate advantage for VOD, MPEG-4 has the additional powerful attributes of interoperability and intrinsic interactivity. In addition, MPEG-4 has wisely included an interface to content protection systems called the Intellectual Property Management and Protection (IPMP) protocol.

One of the most important lessons we learned from the 1990s digital television and cable modem experiences is the critical importance of open standards. Even in an increasingly software oriented world, open standards are the best method for rapid adoption of new products and technologies, multiple sourcing, and simplifying everyone's life from a product management and customer service perspective. Just as the MPEG-2 and DOCSIS standards have aided the entire constituency chain, from content providers to manufacturers to cable operators to consumers, MPEG-4 will perform a

similar role in the IP streaming media arena. In fact, standardization of MPEG-4 is even more crucial since IP is the basis for the first truly global network.

There are three predominant proprietary streaming formats today, RealNetworks, Windows Media Technology and QuickTime, all of which incorporate proprietary video codecs. These formats have improved significantly in the last year and are proving to be effective in enabling streaming media businesses to get launched. Broad proliferation of high quality streaming, however, will require widespread adoption of a standard such as MPEG-4. The current "multiple proprietary format" situation makes for a painful and needlessly complex situation, especially for content providers and consumers, forcing questions such as which format(s) to use, which player(s) to download, etc. While MPEG-4 alone is not a panacea, its adoption will help move the industry toward a common interoperable platform for streaming media. Just as the momentum of MPEG-2 superseded the early proprietary digital television compression techniques, today's proprietary streaming codecs will be pushed aside by MPEG-4.

While MPEG-2 is confined to audiovisual coding and synchronization, forcing a patchwork of proprietary interactive extensions, MPEG-4 is a far broader object-oriented coding standard. MPEG-4 promises to be the holy grail of interactive media, offering content providers the ability to combine real video and audio with synthetic audiovisual objects and text, allowing, for the first time, a unified and integrated standards-based framework for creating super rich interactive content.

With digital television, MPEG-2 video was important but insufficient to ensure full interoperability, and various other standards came into play. Similarly, while MPEG-4 will become the new core of broadband

streaming media, other new and emerging standards will be used in conjunction with MPEG-4 to create an end-to-end system.

- **Audio:** While MPEG-4 allows a variety of audio codecs to be used, the Advanced Audio Coding (AAC) standard is the preferred one for high quality audio. The roots of AAC began several years ago when it was proven within the MPEG-2 standards committee that a surround sound audio system burdened by backward compatibility to MPEG-2 Layer 2 audio (Musicam) would be inferior to a state-of-the-art system without such a burden. Thus began the Non Backward Compatible (NBC) audio mode of MPEG-2, which later morphed into AAC. Offering near CD quality stereo at 96-128 kbps, significantly more efficient than MP3 or Dolby® Digital AC-3, AAC is destined to play an important role in the future of streaming media.
- **Transport:** The MPEG-2 systems layer has become the dominant transport layer for digital television, encapsulating elementary video and audio streams into 188 byte packets, including a 4 byte header and time stamps for synchronization.

Transmission Control Protocol/Internet Protocol (TCP/IP) is the universal transport/network layer for Internet traffic. Related to TCP/IP are various other protocols pertaining to streaming. User Datagram Protocol (UDP) is widely used instead of TCP for streaming traffic. UDP is a more appropriate transport layer than TCP for streaming as it is more of

a broadcast oriented protocol and has lower overhead. Unlike TCP, UDP does not incorporate packet sequencing or retransmissions (through acknowledgments), and doesn't require a set up handshake between the server and client.

- Real Time Protocol (RTP) allows packetized transport of audio, video and data streams. RTP is used in the ITU H.323 standard for video conferencing over the Internet, and is becoming the de facto standard for IP multimedia transport (i.e., of MPEG-1, MPEG-2, MPEG-4, etc.). RTP is used in conjunction with Real Time Control Protocol (RTCP). Together, these two protocols provide such functions as detecting out of order packets, estimation of packet loss and jitter, payload identification, Quality Of Service (QoS) feedback, and multiplexing of media streams.
- Real Time Streaming Protocol (RTSP), used in conjunction with Session Description Protocol (SDP), provides session based VCR-like functions such as fast forward, rewind and pause, as well as metadata information about the program itself. Therefore, RTSP is the protocol which fundamentally allows VOD over IP. RTSP was built on top of Hyper Text Transfer Protocol (HTTP), using the HTTP header definitions and status codes, while adding specific RTSP oriented headers and codes to this list.
- Hyper Text Transfer Protocol (HTTP), the basic protocol used by Web servers for sending

Internet traffic, can also be used to transmit video files to end users. Typically called “progressive download,” HTTP video can be thought of as pseudo streaming since it's a form of file download in which the end user can start viewing the video before the entire file has been downloaded. While true streaming has various long term performance advantages, HTTP progressive download has the interim advantages of allowing the use of pre-existing Web servers (instead of dedicated streaming servers), and also getting through corporate firewalls.

The following chart shows the major technological elements of a digital television system alongside the analogous elements of an end-to-end broadband streaming media system.

Technology Element	Digital TV	Streaming Media
Video Coding	MPEG-2	MPEG-4
Audio Coding	Dolby®Digital (AC-3); Musicam	AAC
Transport/ Systemization	MPEG-2	MPEG-4, IP, UDP, RTP, RTSP, MPEG-2
Interface to Content Protection System (DRM)	MPEG-2 PSI (Program Specific Information)	MPEG-4 IPMP
Digital Rights Management/ Conditional Access	DigiCipher® II, PowerKey®	Motorola/ Aerocast™ DRM, Microsoft DRM
Content Guide	TV Guide, TV Gateway, Native guides	Internet publishing directory with search engine and branded portal

Secure Digital Rights Management (DRM)

Good quality low bit rate innovations such as MPEG-4 present both an opportunity and a dilemma for copyright holders and content distributors. While the notion of a new global distribution channel for video content is inherently attractive, conventional wisdom is that the Internet is an insecure medium. With MP3 Internet file sharing services such as Napster giving content rights holders a definite scare, secure digital rights management (DRM) becomes a critical element for any viable broadband streaming video solution. DRM involves establishing and enforcing rules for the distribution and use of content. Simply put, subscription and PPV oriented content providers cannot risk putting their video assets on the Internet in a way that financially jeopardizes their other revenue sources and distribution partners. And many new content providers will want to utilize DRM in order to monetize their media assets.

Digital rights management is a broader category of content protection than the cable industry's traditional notion of transport-oriented conditional access, with the primary addition being the concept of persistent rights management, defined as the ability to track and protect copying and re-distribution of content from source(s) to destination(s).

A streaming media DRM system should incorporate the following elements:

- Authentication: securely verifying that clients and servers attempting to participate in the streaming process are indeed who they claim to be
- Privacy: rendering the content, as well as other sensitive information, unintelligible to parties which don't have a valid decryption key

- Message Integrity: ensuring, by cryptographic techniques associated with a message, that the message, when delivered to a recipient, was not modified or tampered with
- Conditional Access: controlling the ability of individual subscribers to view or listen to content on a subscription, PPV, or other restricted basis
- Persistent Copy Protection: applying rules which dictate who can do what to content with respect to copying or re-distribution, as well as techniques such as watermarking and fingerprinting which can help trace the source of the content

In designing a state-of-the-art streaming media DRM system, proven and existing standards should be employed such as the recently adopted Advanced Encryption Standard (AES), represented by the powerful Rijndael encryption algorithm. A combination of public key and symmetric key techniques should be used. Public (asymmetric) key cryptography involves a public key, which is generally known, used in conjunction with a different private (secret) key. With symmetric key cryptography, the encryption and decryption key are the same, and this key must be kept private and secret. In general, Public Key Infrastructure (PKI) techniques should be used for cryptographic applications that are less time-sensitive, such as initial set-up and authentication, while the latest symmetric key techniques should be utilized for time-sensitive applications such as payload encryption.

Major system features should include subscription, pay-per-view, pay-by-time, unicast, multicast, parental rating controls and location-based access, as well as the ability to control copying and re-distribution.

Intelligent Content Distribution

Delivery of digital television content is a mature and fairly straightforward process including satellite links for backhauls and distribution, microwave links, signal collection in headends, and modulation and upconversion for re-distribution through the cable plant. Additional intelligence is required in order to repackage individual digital programs from multiple transponders into a single cable multiplex. In the MPEG-2 digital TV world, with the emergence of VOD, PVRs and user-based preference schemes, content delivery will become incrementally more refined.

But with IP streaming media, intelligent content distribution takes on a whole new meaning, heralding a hyper dynamic system of global, national, regional and local networked VOD in which content will be cached where and when it is most demanded, while preserving for content providers the flexibility to broadly distribute content if that is what they desire and can justify economically. A key issue for IP streaming media is edge caching (storing and serving content from a point as electronically close to the end viewer as possible). But mere edge caching without intelligent content distribution is tantamount to distributing magazines to newsstands around the world without regard to the particular languages, tastes or preferences within each localized vicinity surrounding the newsstand.

The following concepts exemplify intelligent content distribution in the context of a next generation broadband streaming media system:

- Pre-positioning: moving a streaming content file in its entirety to a pre-selected set or subset of edge servers. This method will be used by content providers which are reasonably confident that one or more subscribers in each edge server vicinity will want to view the content, and are therefore willing to invest upfront for broad distribution and storage.
- Partial pre-positioning: moving a portion of a streaming content file to a pre-selected set or subset of edge servers. This method is more conservative than full pre-positioning and will still mitigate the phenomenon whereby the first viewer of a piece of content is penalized by drawing the stream all the way from the origin server or from a non-optimal edge server. Various tradeoffs can be contemplated between how large of a portion of the content object to pre-position at the edge(s) relative to the projected network latency issues surrounding streaming the remaining non-pre-positioned portion of the content object.
- Demand based caching: For broad libraries of content, as well as niche, local or individual self-published content, it may make more sense for the object to remain at the origin server site until a subscriber requests it from the portal or directory. Then, the object would begin caching at the preferred edge server, and the initial requestor would have an acquisition penalty. Subsequent viewers from the same edge server would be able to view the content with high performance and without delay.
- Edge server selection: With a network of edge servers, the process and method by which selection of the best edge server at

any given time is a fundamental part of intelligent content distribution. Various techniques exist and are being refined, some of which are significantly more precise than others in their ability to select the electronically closest edge server.

Additional factors pertain to intelligent content distribution such as how long to store an object at an edge server (e.g., absolute time period, pre-determined time period from the last time the object was requested, etc.), as well as the integration of the intelligent content distribution system with the DRM system.

Proliferation of Broadband IP Devices

There are now over 5 million broadband Internet subscribers in the United States, the majority of which are cable modem users. This figure is expected to grow to at least 30 million households by 2005. In addition, cable modems are increasingly being integrated into digital set-top boxes. Therefore, over the next few years, broadband Internet access will become not just a PC-centric phenomenon but also a TV-centric enhancement, and streaming media will be the driving force which accelerates and sustains this trend.

Beginning with the advanced set-tops being introduced in 2001, increasingly more powerful microprocessors will be incorporated allowing many additional functions and applications to be implemented in software. In addition, as standards such as MPEG-4 take hold, cost-effective dedicated hardware devices such as MPEG-4 circuits will be introduced.

Cost Effective and User Friendly Software and Hardware Tools Leading to an Explosion of Internet Content

Perhaps the greatest promise of broadband streaming media is that the finite shelf space shackles of the television world will finally be torn down, allowing virtually anyone to be a content provider. Of course business models and the development of audiences and viewers will still ultimately determine success, but the growth in content availability enabled by broadband streaming media will be mindboggling and exponential relative to the impressive increase in content offerings which occurred due to the move from analog to digital television.

With broadband streaming media, each and every content genre (sports, education, travel, culture, music, kids, health, cooking, weather, just to name a few) will become fully represented in a way that television cannot replicate. The concept of narrowcasting and the availability of niche and special interest content will be taken to a whole new level. The local content enabled by streaming media will give operators a powerful new content weapon in the race against DTH providers, upstaging the limited “local into local” capability provided to DTH subscribers.

This is not to say that streaming will replace broadcast television; it will not. Television will remain the superior technology for mainstream popular content for many years to come. Instead, streaming will complement television in a synergistic way that is only starting to be conceived.

Helping this trend along is the increasingly cost effective and user-friendly availability of software and hardware tools. For example, the new generation of Digital Video (DV) camcorders is not only starting to become affordable, but also is just as easy to use as yesterday’s analog camcorders.

These DV camcorders have standard IEEE 1394 (Firewire/i.Link) interfaces for PCs, allowing the captured motion video to easily be transferred to hard drives in PCs, where the content can be edited and manipulated with software which is becoming less expensive, more user friendly and more powerful every day. Finally, software based encoding, streaming, distribution and directory tools will allow this content to be made available to broad or narrow audiences on a global basis.

Conclusion

A combination of new technologies and economic conditions is materializing which will usher in a new era of broadband streaming media. Broadband streaming media is unprecedented in its ability to deliver on demand (and live) content of every shape and form. It is important and timely for operators to weigh in on important issues such as the MPEG-4 and RTP/RTSP standards. The emerging MPEG-4 standard, secure digital rights management, and intelligent content distribution are among the most prominent of the technologies that will form the foundation for this exciting new media category. Enormous opportunities await the content providers, technology companies and operators that embrace and tap into the future of streaming media.

Contact Information:

Marc L. Tayer
Senior Vice President, Business
Development
Aerocast, Inc.
5744 Pacific Center Boulevard
San Diego, CA 92121
(858) 404-6211
mtayer@aerocast.com

THE ROLE OF SERVICE PROVIDERS IN HOME NETWORKING INCREMAL REVENUE OPPORTUNITIES FOR MSOS

Don Apruzzese
ShareWave, Inc.

Abstract

With the rise of multi-PC households and broadband subscribers, high-performance home networks are required to easily extend the broadband experience throughout the home to numerous devices—set-tops, mobile web pads, and other information appliances—while preserving quality of service (QoS). Therein lies significant revenue opportunities for MSOs, not just in adding premium services, but also in delivering them to additional devices connected to the broadband access point on a home network. The solution, however, must be high-performance, easy to install, affordable, and scalable. This paper addresses the requirements for distributing multimedia and digital content throughout the home, and the associated benefits to service providers—including consumer self-installation, limiting device redundancy, and incremental revenue opportunities. Target audience: Service providers, OEMs, press, analysts

THE ROLE OF SERVICE PROVIDERS IN HOME NETWORKING

“In three to five years, home networks will have moved off the pages of science magazines and into millions of US homes. And millions of people, seeing a home network as a huge opportunity to...obtain valuable new services, will embrace the technology with enthusiasm.”

—The McKinsey Quarterly, 2001 Number 2

Introduction

The increasing penetration of broadband technology into the home is a key driver for home networks, which extend the broadband experience to multiple computers and digital devices anywhere in the home.

In addition to benefiting broadband subscribers, home networks enable broadband service providers to increase revenues, decrease installation costs, and accelerate the deployment of new broadband services, such as video on-demand (VoD), interactive television, and IP-based media services.

This translates into significant competitive advantage for the MSO: capturing and retaining more subscribers while opening up new value-add revenue streams.

To take advantage of these incremental revenue opportunities, it's in the MSO's best interest to aggressively encourage the deployment of home networking technology.

Because these incremental revenue generators are multimedia-focused, they require home network technology that seamlessly transmits high-fidelity digital content and supports true quality of service (QoS).

ShareWave's Vision

The deployment of broadband content and services, distributed throughout a home network, is changing how people interact with digital content. ShareWave's vision is to enable *wireless* home networks that provide consumers with access to a variety of digital content—entertainment, information, and communications—when and where they want it. This ubiquity makes consumer interaction with digital content more enjoyable, convenient, and affordable. More importantly, it raises the subscriber's perceived value of the content, leading to incremental revenue opportunities for the MSO.

Emergence of New Devices and Applications

The proliferation of new digital devices, content, and services further drives the broadband and home network markets. Beyond traditional PCs and laptops, consumers now have a choice of many exciting new devices:

- Set-top boxes
- Integrated personal video recorders (PVRs)
- Cable modems
- Residential gateways
- Digital audio and video jukeboxes
- MP3 players
- Mobile web pads
- Advanced game consoles
- Other internet appliances

These devices facilitate—and increase—consumer interaction with digital content and broadband services.

Some application examples:

- Listening to MP3 audio files through the home stereo system, instead of PC speakers

- Internet appliances, such as a kitchen pad, that allow for convenient access to broadband content appropriate for activities conducted in a specific room.
- One PVR providing recorded content to multiple TV sets throughout the home, versus restricting it to one TV.

Without the deployment of home networks to connect these devices, consumer consumption of these new content offerings and services will not achieve its full potential.

Delivery of Broadband Services

The deployment of broadband services fits hand-in-hand with home networks. Both are required for consumers to gain full benefit.

Increasingly, consumers want more than just simple file and print sharing. Studies show that subscribers with cable modem service are significantly more likely to access multimedia content than dial-up subscribers. Consumers are interested in entertainment and interacting with digital content—music, video and games.

The ability to distribute this content to multiple locations throughout the house makes broadband services much more appealing. Anyone in the family can interact anywhere, anytime, from any device.

This creates an ever-important value chain: devices are being developed to support new and valuable broadband content and services. The more accessible the broadband connection is to these devices, the more the subscriber will use and value it. And the more they will pay for it.

Therefore, home networks add significant value to services such as:

- Video on-demand (VoD)
- Interactive television
- IP-based media services
- Online gaming
- Broadband Internet sharing
- Voice over IP (VoIP)

Looking ahead, home networks also provide MSOs the opportunity to extend their services beyond the home—to the automobile, for example, where average subscribers spend up to 3 hours a day while commuting. The subscriber could listen to personalized music or news that was downloaded from the home server or broadband connection while the automobile was parked in the garage the night before.

Service Provider Benefits

- Easy, low-cost installation

Home networking enables service providers to initially reduce the time to install new broadband services by eliminating in-home wiring, thereby reducing the cost of installation.

Eventually, home networks will eliminate expensive truck-rolls entirely by supporting easy subscriber self-installation.

- Acceleration of service deployment

New broadband service deployments are constrained by the MSO's inability to install the customer equipment fast enough.

Reducing installation time, and subsequently moving to self-installation due to the home network, will accelerate subscriber penetration rates and thus revenue from those new services.

- Support for retail equipment model

Because a new device can be added to the wireless network easily, consumers can purchase new devices at retail stores, self-install, and order the new service(s) without the need for a truck-roll. This improves marketing efficiencies and inventory carrying costs.

- Incremental and enhanced revenue

With home networks, MSOs have the unique opportunity to implement additional revenue streams, including:

- Network maintenance and management, including software updates and services
- Tiered IP service levels, based on bandwidth requirements and number of devices
- Tiered service packages for premium video, VOD, audio and interactive gaming
- Personalized services, based on device usage and content consumption

- Reduced device redundancy

A home network supports the concept of a single, full-featured set-top box, cable modem, residential gateway, or media server to deliver broadband services to several less-expensive devices, thereby eliminating the cost of additional boxes. For example:

- One cable modem supporting multiple PCs or Internet devices.
- Distributing the cable modem connectivity within an advanced digital STB to one or more PCs located in other rooms of the home.
- Distributing the capabilities of one fully featured STB (CM & PVR) to one or more "thin" STBs attached to other TVs or devices.

- Differentiated competitive service offerings

By deploying a home network to distribute content and services to multiple devices, the MSO is able to create a differentiated service offering versus its competition. This lends itself to improved new subscriber acquisition.

- Subscriber retention

As the home network becomes an integral part of the service offering, the MSO is effectively creating a barrier to entry for the competition. Subscribers will become less likely to switch service providers because of their dependence upon the multiple services offered by the MSO. Customer retention would improve in the face of competition.

The home network could offer the MSO greater visibility into the subscriber content and service usage patterns on every device in the home. This information is invaluable in developing new service offerings and improving customer satisfaction.

- Customer service

Remote diagnosis of networked devices from the headend allows for better customer service, troubleshooting, and reduced cost by eliminating truck rolls when possible.

Multimedia Requirements

Although broadband Internet access offers significant revenue opportunities for MSOs in the short-term, entertainment and voice services—whether IP- or traditional MPEG transport-based—offer the highest revenue opportunity over time.

Therefore, any home network must support multimedia and address the requirements for transmitting time-dependent, digital content.

- Breadth of applications

- Connectivity for applications on a wide-range of devices including STBs, PCs, consumer electronics, mobile pads, cable modems, Internet appliances, and more.

- Support for quality of service (QoS)

- Flawless support for multiple isochronous media streams (video, voice, audio); Greater throughput for transmitting high fidelity digital content, such as MPEG-2 video and CD-quality audio; Efficient bandwidth allocation and usage to support multiple simultaneous content streams; Predictable delivery of time-dependent content.

- Ease of use

- Open enrollment for easily adding devices; Peer-to-peer communications among multiple devices; Efficient bandwidth usage; Master redundancy; No single point of failure.

- Reliability

- Avoidance of in-band interference from household appliances, such as microwave ovens and cordless phones; Sufficient range and fade margin for coverage throughout the home; Use of Forward Error Correction (FEC) to minimize bit error rate (BER) and support time-dependent content without retransmission.

- Scalability

- Obsolescence protection to preserve device investment
- Foundation for future features, devices, and services.

Other Home Networking Issues

- Price/performance is a key driver for home networking. Currently transmitting up to 11 Mbps.
- Home networks will be heterogeneous environments, incorporating WANs, LANs, and PANs.
- 802.11b is emerging as the wireless standard for home networks, but lacks QoS and multimedia support

Close

It is in the best interest of all MSOs to rapidly deploy home networking technology. Herein lies a unique opportunity for MSOs to maximize their incremental revenue opportunities, increase customer retention and fend off the increasing competitive threat. The technology is available today.

Author Information

Don Apruzzese
Director of Business Development,
ShareWave, Inc.

Don Apruzzese serves as Director of Business Development for ShareWave, Inc. In this capacity, Apruzzese works closely with leading OEM suppliers to the telecommunications industry to integrate ShareWave's wireless networking technology into next generation digital set-tops and cable and xDSL modems. Prior to ShareWave, Apruzzese served as Director, Product Marketing Management, Broadband Internet Services of MediaOne. In this position, Apruzzese was responsible for overseeing the company's enterprise-wide marketing effort for its broadband Internet products. Apruzzese was appointed Director, Product Marketing Management for Continental Cablevision, Inc., in March of 1996, prior to its merger into Media Group. Continental became MediaOne in May of 1997.

Prior to joining Continental Cablevision, Apruzzese held several key marketing positions with Ford Motor Company and Braun, Inc. Apruzzese received his B.F.A. in Film Production from New York University, Tisch School of the Arts, and his MBA from The Wharton School, University of Pennsylvania.

916-939-9400 x3201
don@sharewave.com

Company Information

ShareWave provides semiconductor technology for multimedia-capable wireless home networks. ShareWave delivers the industry's optimal solution for wirelessly transmitting the full range of high-fidelity multimedia content—including MPEG-2 video and CD-quality audio—to complement broadband delivery into the home.

The company is privately held with its headquarters in El Dorado Hills, CA. ShareWave has received funding from APV Technology Partners, Cisco Systems, Intel Corporation, KLM Capital, Kyushu Matsushita Electric Co., Ltd., Microsoft Corporation, Philips Electronics NV, SBC Communications, Inc., SOFTBANK Technology Ventures, Vulcan Ventures, Inc., and other public and private investors.

For more information, visit the company's website at www.sharewave.com.

THERMAL DESIGN AND TEST REQUIREMENTS FOR OUTSIDE PLANT CABLE TELECOMMUNICATIONS EQUIPMENT

Al Marshall, P.E.
Philips Broadband Networks

Abstract

Shrinking thermal margins, driven by sophisticated but thermally sensitive components, greater heat dissipation and industry requirements to operate in more challenging environments, threaten the reliability of outside plant equipment in advanced broadband networks. Simply specifying an ambient temperature of -40°C to +60°C is no longer adequate for outside plant equipment because the effects of wind speed, installation type and solar radiation can eliminate the thermal margin.

THE PROBLEM

More Capability, Higher Temperatures

The demand for advanced, bandwidth hungry telecommunications and entertainment services will force system operators to add more functionality to their networks. Equipment power consumption and associated heat dissipation will increase to support these new services. The additional heat dissipation will drive the operating temperatures of lasers and other temperature sensitive components up, particularly in outside plant (OSP) equipment where active cooling is unattractive.

The test methods commonly used to exercise OSP equipment over temperature do not replicate the real world environment accurately or even conservatively. This masks potential field problems caused by inadequate or negative thermal margin. Small or negative thermal margins reduce system availability.

Thermal Margin

The performance of electrical components and other materials used in optical nodes, RF amplifiers and other OSP equipment varies over temperature. There is a tem-

perature range for each component in a piece of OSP equipment over which that component will provide adequate performance for acceptable system operation. The maximum and minimum temperatures in this range define the application specific temperature limits for each component.

Component temperatures in fielded OSP equipment are affected by ambient air temperature, solar loading, altitude, wind speed and wind direction. The materials and design of enclosures and the installation of the OSP equipment in enclosures further affect the temperatures of components housed in vaults, cabinets and pedestals. Component heat dissipation, the dissipation of neighboring components and the total dissipation of the entire piece of OSP equipment are other significant factors that influence component operating temperature.

The difference between the operating temperature of a component and its application specific temperature limit is the thermal margin. The thermal margin is negative if the component operating temperature exceeds either of the application specific temperature limits.

Improvements in cable telecommunications technology will push component operating temperatures up while maintaining or reducing the maximum application specific temperature limit.

Increasing Heat Dissipation

New services, improved transmission schemes and more capable network management will be implemented in HFC networks. These advancements will add hardware and heat to network equipment.

Incumbent Local Exchange Carriers (ILECs) have already begun to see the effects of network advancements. A conventional

POTS (Plain Old Telephone Service) line consumes about 2 watts of central office power. ADSL lines consume between 6 and 8 watts. A central office that consumed 2000 to 3500 amps of electricity ten years ago may use 10,000 amps today. An ILEC that had 10 central offices in a medium-sized city in 1990 may now have 10 central offices, 150 vaults and 1,000 cabinets¹. Most power provided to telecommunications systems becomes waste heat. Since the cable telecommunications industry plans to provide similar services, a similar growth in heat dissipation is a likely outcome.

The computer industry is another example of increasing heat dissipation. Dissipation of computer CPU's has increased significantly with processing capability and data rates. Figure 1 and Figure 2 illustrate this point for Intel and AMD processors².

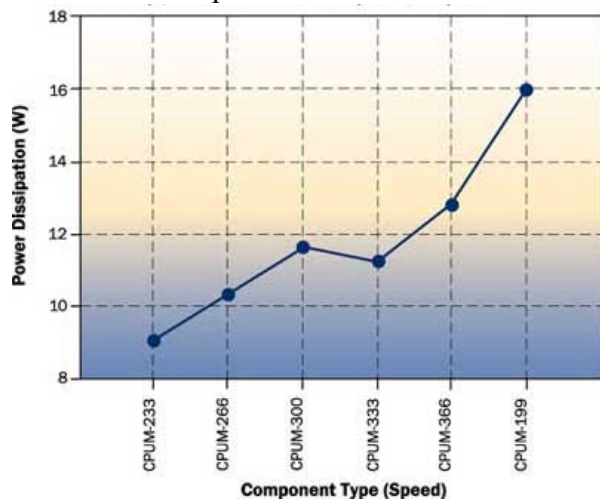


Figure 1 - Intel CPU Heat Dissipation vs. Clock Speed

The combined effects of new services, improved functionality and increased capacity will bring an increase in the heat dissipation of HFC network systems and their components.

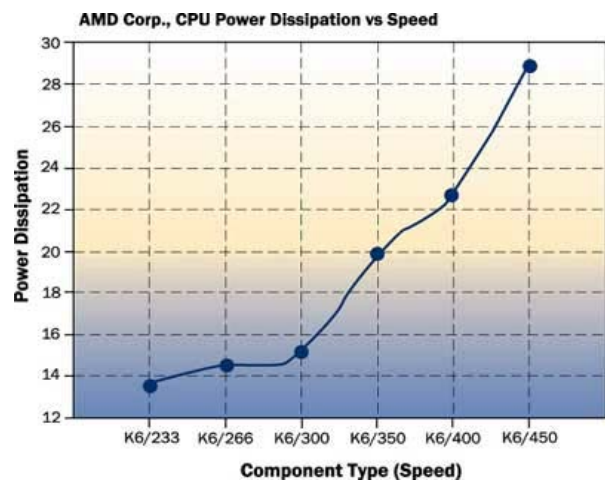


Figure 2 – AMD CPU Heat Dissipation vs. Clock Speed

Thermally Sensitive Components

RF amplifiers tested outside the normal operating temperature range for a short duration usually exhibit graceful performance degradation. The components being used or being considered for use in HFC networks today do not always degrade gracefully.

Optoelectronic components are necessary in the outside plant. Lasers are available with maximum case temperature ratings of 65°C to 85°C depending on the intended application. Lasers exhibit a sharp, ungraceful reduction in optical output power as case temperature increases past the maximum temperature limit. Lasers designed for use on the ITU wavelength grid also suffer from unacceptable wavelength drift at high temperatures. Both failure modes result in an unusable signal path.

Digital devices exhibit a variety of failure modes at extreme temperatures. The timing of critical high-speed signals may become skewed causing a breakdown in device to device communication. Random thermal noise can obscure low-level signals. Some devices will “lock-up” at temperatures outside the specification limits. These failures may cause degraded system performance or they can cause complete system failure.

The continued use of optoelectronics and the addition of digital electronics will maintain or lower the maximum application specific temperature limits for components used in OSP equipment.

First Order System Thermal Model

The remainder of this paper will focus on fiber optic nodes mounted in pedestal enclosures. The heat transfer mechanisms, design requirements and test methods presented are applicable to RF amplifiers and other OSP equipment.

Component operating temperatures in OSP equipment are a function of many variables. Figure 3 is a first order model that illustrates the relationships between key parameters for a node installed in a pedestal enclosure.

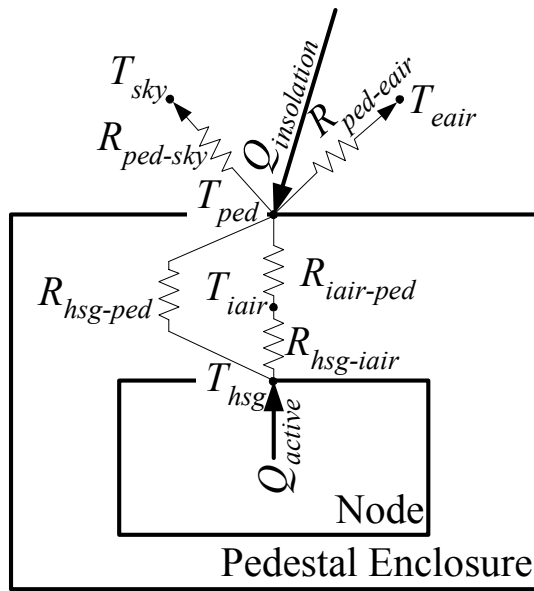


Figure 3 – First Order Thermal Model of a Node Installed in a Pedestal Enclosure with Local Wind Speed Equal to Zero

Q_{active}	Heat Dissipated by Node
$R_{hsg-iair}$	Thermal Resistance from Node Housing to the Internal Air (Convection)
$R_{hsg-ped}$	Thermal Resistance from Node Housing to Pedestal Enclosure (Radiation)

$R_{iair-ped}$	Thermal Resistance from Internal Air to Pedestal Enclosure (Convection)
$R_{ped-eair}$	Thermal Resistance from Pedestal Enclosure to External Air (Convection)
$R_{ped-sky}$	Thermal Resistance from Pedestal Enclosure to Sky (Radiation)
T_{hsg}	Average Node Housing Temperature
T_{iair}	Pedestal Internal Air Temperature
T_{ped}	Pedestal Enclosure Temperature
T_{eair}	Ambient (External) Air Temperature
T_{sky}	Sky Temperature
$Q_{insolation}$	Heat Transferred to the Pedestal Enclosure by Direct Solar Radiation

Equation 1 governs the temperature rise across the resistors in Figure 3.

$$\Delta T = Q R_{thermal} \quad (1)$$

Q is the heat passing through each resistor. The sky and external air temperatures are not controllable. As Q increases, the temperature difference across each resistor increases.

The values of $R_{thermal}$ depend on the mode of heat transfer. Heat transfers by conduction, convection and radiation. Conduction is not a significant mode of heat transfer in pedestals. Convection can be free (natural) or forced. Thermal resistance for convection is determined from equation 2.

$$R_{Thermal} = \frac{1}{hA} \quad (2)$$

h is the convective heat transfer coefficient. A is the surface area of the node or other OSP equipment.

The value of h is calculated differently for free and forced convection. h under forced convection depends on air velocity, length in the direction of flow and several fluid properties. h under natural convection depends on surface orientation with respect to gravity, the temperature difference between the surface and the surrounding fluid and several fluid properties. Analytical estimation of h for real systems is complex and is usually done numerical techniques such as Computational Fluid Dynamics (CFD).

Heat transfer by radiation follows equation 3³.

$$Q = \sigma f e A (T_{hsg}^4 - T_{ped}^4) \quad (3)$$

σ is the Stefan-Boltzmann constant, f is the view factor, e is the emissivity factor and A is the surface area of the OSP equipment. For a first order approximation of a node in a pedestal, e is approximately equal to the emissivity of the node surface and f is approximately equal to one. CFD codes optimized for electronics cooling applications include a radiation analysis capability for better estimates of heat transfer by radiation.

The node housing temperature is the sum of the temperature differences across the resistors added to the external air/sky temperature.

Internal component operating temperatures (e.g. return transmitter lasers) are dependent on the node housing temperature. Components will approach and eventually exceed their application specific temperature limits if the heat dissipation of cable telecommunications equipment follows the trends in the computer and telephone industries. The consequences of this trend are shown in Figure 4.

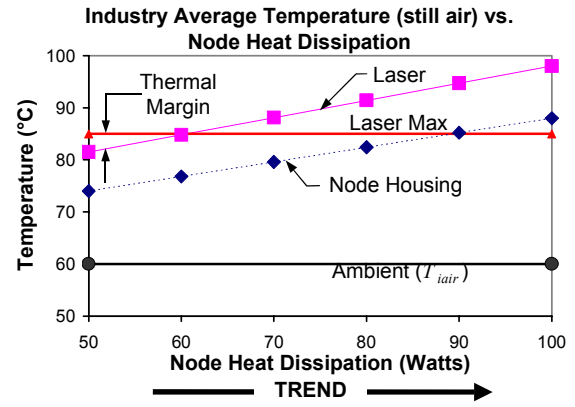


Figure 4 – Industry Average Node Temperature vs. Node Heat Dissipation in a Still Air Ambient Environment (Pedestal)

The Real World vs. Environmental Chambers

Figure 4 depicts node housing and laser temperatures as a function of node heat dissipation in still air. Still air conditions can exist in pedestal enclosures, cabinets and underground vaults. Most high and low temperature OSP equipment tests are run in environmental chambers. Environmental chambers employ moving air to keep the chamber temperature uniform. $R_{hsg-iair}$ for a node in a chamber will be much lower than $R_{hsg-iair}$ for a node operating in still air. The reduction of $R_{hsg-iair}$ will lower the node temperature relative to the surrounding air. Figure 5 shows how testing a node in a chamber can hide the negative thermal margin that will exist when the same node is operating in a pedestal in the outside plant.

The laser that appeared to have adequate thermal margin in the environmental chamber (Figure 5) only has adequate thermal margin in a fielded pedestal if its node dissipates less than 60 Watts of heat (Figure 4).

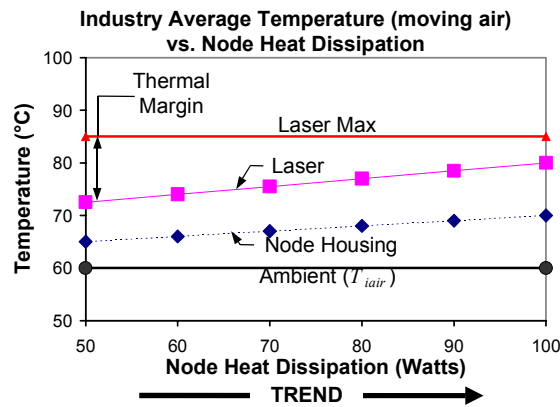


Figure 5 - Industry Average Node Temperature vs. Node Heat Dissipation in a Moving Air Ambient Environment (Environmental Chamber)

Increasing heat loads, constant or declining application specific component temperature limits and non-conservative test methods will result in negative thermal margins in future generations of OSP equipment. Design requirements that describe and test methods that reproduce the OSP environment more accurately are necessary to ensure that equipment delivered to the field is thermally robust.

NODE HEAT TRANSFER IN A PEDESTAL

An understanding of pedestal heat transfer mechanisms is necessary to write design requirements and develop tests for a pedestal environment. Figure 3 models the effects of these mechanisms. The resistor values depend on many variables.

Natural convection and radiation cool a node operating in a pedestal when the external air is still. Natural convection moves heat from the node to the surrounding air. Heat from the node warms the adjacent air. The warm air rises to the top of the pedestal enclosure and moves from the center of the enclosure to the sides. The air cools as it comes in contact with the relatively cool sides of the enclosure. The cool air sinks to the bottom of the enclosure where it waits to be “sucked” past the node. The flow lines in Figure 6 illustrate natural convection airflow in a pedestal enclosure.

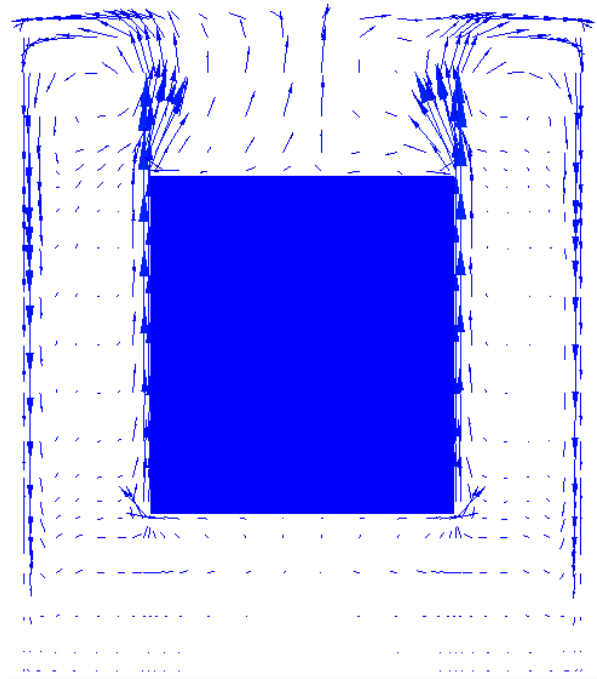


Figure 6 – Natural Convection Airflow around a simplified Node in a Pedestal without Solar Loading

Heat transfers by radiation from the node to the pedestal enclosure and from the node to the ground at the bottom of the enclosure when the node is at a higher temperature than the enclosure or the ground. Radiation heat transfer between the node and the pedestal enclosure can vary significantly depending on the emissivities of the surfaces and the temperature differences between the surfaces. If the interior surface temperature of the pedestal walls becomes greater than the temperature of the node surface, heat will transfer by radiation from the pedestal enclosure to the node. This could happen with certain pedestals under solar loading.

Figure 6 depicts a pedestal in still air with no solar loading. Solar loading changes the internal airflow pattern. Direct solar load will irradiate two vertical walls and the top of the pedestal depending on the orientation of the pedestal and the position of the sun. Other walls may receive solar radiation reflected from nearby walls, sidewalks and other terrain features. The warm air rising from the node

must pass the non-irradiated walls for cooling if the internal temperature of the irradiated walls is equal to or greater than the temperature of the node. The asymmetrical and more restricted airflow will cause $R_{hsg-air}$ to increase. The solar load on the pedestal will increase T_{ped} . The node housing temperature, T_{hsg} will increase. The magnitude of the increase will depend on the materials used in the pedestal, the orientation of the sun, the pedestal size, the node size and the location of the node within the pedestal.

Wind will pass through the ventilation slots in the pedestal and mix with the internal air. The wind speed inside the pedestal will be less than the outside wind speed. The reduction in wind speed depends on the pedestal orientation, size, vent geometry and vent cleanliness. Local obstructions such as landscaping, buildings and fences also affect internal wind speed. Increasing internal wind speed will lower $R_{hsg-air}$ and reduce $R_{iair-ped}$. A new parallel resistance, $R_{iair-eair}$, reduces the overall resistance between the housing and the external ambient air. See Figure 7.

Wind will result in lower internal air and node housing temperatures than still air. Still air should be assumed for design purposes because it represents a conservative but occasionally real field condition.

One consequence of pedestal fluid mechanics is higher air temperatures at the top of the pedestal than at the bottom. Temperature differences of up to 20°C have been measured between the top and the bottom of fielded pedestals.

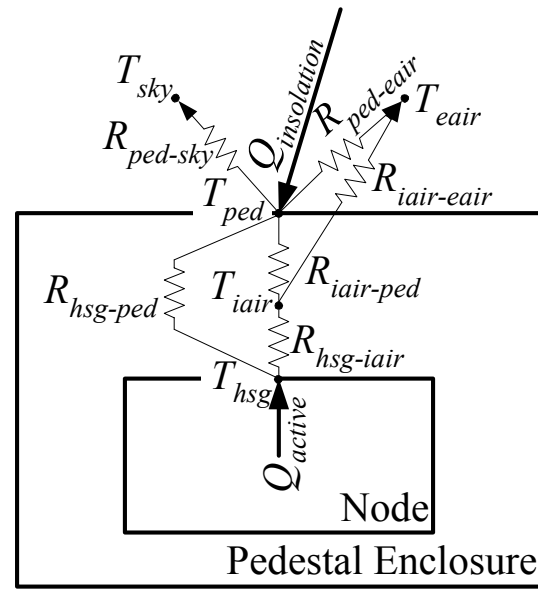


Figure 7 - First Order Thermal Model of a Node Installed in a Pedestal Enclosure with Local Wind Speed Greater than Zero

DESIGN REQUIREMENTS

Thermal design requirements for any electronic product must address the three primary mechanisms of heat transfer; conduction, convection and radiation. Heat transfers from a heat source to a heatsink. Heatsink temperatures and parameters that determine the thermal resistance between the heat source and heatsink must be specified.

Conduction from a node to a pedestal or to the surrounding environment is a relatively minor means of heat transfer. Conductive heat paths through the mounting points, coaxial cables and optical fiber cables are long and filled with high resistance materials or interfaces. Assuming that no heat is transferred from the node by conduction is a slightly conservative assumption that is reasonable for design purposes.

Radiation could be significant if the emissivity of the node is high. The effect of radiation could increase or decrease the temperature of the node depending on the temperature and emissivity of the inside surfaces of the pedestal. If radiation is to be assumed as a means of heat transfer, pedestal surface

temperatures, emissivities, dimensions and locations must be specified. The ground temperature and the node location with respect to the ground and the pedestal must be specified also. Radiation can be neglected as a mode of heat transfer if the emissivity of the node is low.

Convection is the primary mode of heat transfer for most electronic equipment operating in pedestals. Still air outside the pedestal should be assumed. Natural convection will be cooling the OSP equipment. The heatsink temperature for natural convection should be the temperature of the isothermal pool of air below the node (see the “Test Guidelines” section). Minimum pedestal size and maximum internal surface temperatures must be specified. Node location and orientation must be defined.

TEST GUIDELINES

Measurement Point for Node Ambient

Suppose a node is tested in a simulated pedestal environment (i.e. still air) where the temperature is 60°C at the top to 40°C at the bottom. If the node meets all of its performance requirements, should it be rated as acceptable for use up to 40°C or 60°C?

The air temperature inside a pedestal should be measured at the bottom of the pedestal at least 3 inches below the bottom of the node. The node draws its cooling air from the air at the bottom of the pedestal. The temperature of the air at the bottom has the strongest influence on the temperature of the node. Test repeatability is another reason for measuring temperature at the bottom of the pedestal. Simulations and tests show that temperature is uniform in the air at the bottom of the pedestal. The location of the ambient temperature sensors is not critical as long as they are mounted below the thermal boundary layer under the node.

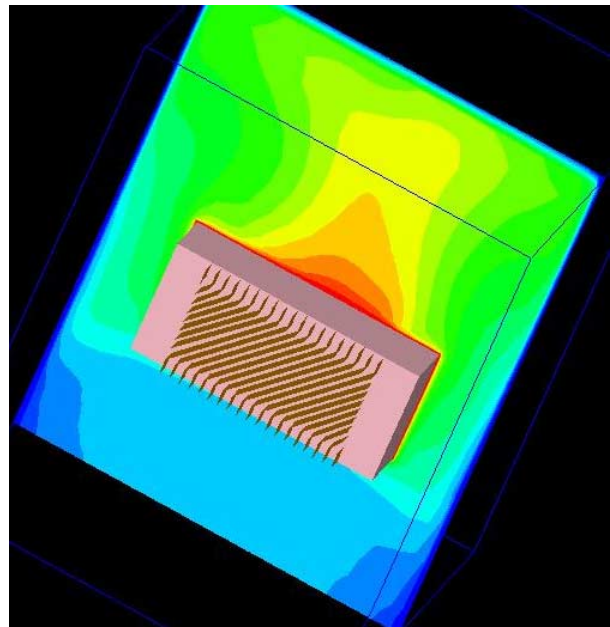


Figure 8 – Predicted Air Temperature Around a Node in a Pedestal in Still Air Without Solar Loading

The contour lines in Figure 8 are isotherms on a plane that passes through the center of a node in a pedestal. Large temperature gradients exist at the top while the temperature at the bottom is uniform. A slight shift in the location of a sensor at the top of the pedestal could result in a large change in the “ambient” temperature reported by that sensor. A sensor at the bottom could be moved without causing such a variation in the indicated temperature. Figure 8 was produced with CFD software.

Test Procedure

OSP equipment performance testing at high or low temperatures is usually done in environmental chambers. The circulating air in a chamber makes the test environment different than in a fielded pedestal. A correction must be made to the chamber temperature to get the OSP equipment housing up to the temperature it will achieve in a pedestal.

Run a simulated pedestal test to estimate the temperature rise from the ambient below the unit under test (UUT) to the UUT housing.

Perform the test in an environmental chamber that is close to the size of an appropriate pedestal. Instrument the UUT housing with temperature sensors near the mounting points for modules and other heat sources. Several sensors should be placed in the air below the UUT to establish where the thermal boundary layer ends and to verify that the pool of isothermal air has been found. Put at least two sensors in the pool of isothermal air because these sensors will determine the reference for temperature rise calculations. Locate and orient the UUT as it would be in a pedestal.

The heat dissipation of the UUT should be as close as possible to its expected maximum. The expected maximum heat dissipation includes the heat dissipated by all application modules and power converters at their maximum operating temperatures. If the UUT is configured to pass current to other devices in the network, then include the ohmic losses (I^2R) created by the maximum passed current.

Run the simulated pedestal test with the chamber control system (heaters, chillers and fans) turned off and the chamber door closed. Verify heat dissipation throughout the test to ensure that all heat sources are working and that the total heat dissipation is correct. Monitor temperatures often enough to determine when the UUT has thermally stabilized with the isothermal pool of chamber air. Stabilization with the chamber air is achieved when the difference between the average UUT housing temperature and the average air temperature in the isothermal pool reaches a constant value.

$$\Delta T_{Natural} = \bar{T}_{UUT} - \bar{T}_{IsothermalPool} \quad (4)$$

Estimate the absolute housing temperature in a pedestal by adding the maximum rated temperature of the UUT to the temperature difference at stabilization. For example, if the UUT is rated to +60°C and the stabilized temperature difference is 25°C, then the estimated average housing temperature in a pedestal is +85°C.

After the UUT has stabilized and the temperature and power data have been checked, turn on the chamber control system. Adjust the chamber temperature to keep the UUT housing temperature where it was prior to activating the control system. The UUT heat dissipation and average housing temperatures should be the same as they were when the stabilized temperature difference was calculated previously. When the UUT is thermally stabilized with the chamber air, calculate the temperature rise from the chamber ambient to the housing.

$$\Delta T_{Forced} = \bar{T}_{UUT} - \bar{T}_{IsothermalPool} \quad (5)$$

Data Reduction

Calculate the thermal resistance between the UUT and the chamber air.

$$\Theta_{Natural} = \frac{\Delta T_{Natural}}{Q_{UUT}} \quad (6)$$

$$\Theta_{Forced} = \frac{\Delta T_{Forced}}{Q_{UUT}} \quad (7)$$

Calculate an equivalent chamber set temperature, T_{Forced} for any real world pedestal internal temperature, $T_{Natural}$ and UUT heat dissipation.

$$T_{Forced} = Q_{UUT} (\Theta_{Natural} - \Theta_{Forced}) + T_{Natural} \quad (8)$$

Example:

A node consuming 70 true RMS watts of AC power and passing 15 amps of AC current is tested as described above. The measured temperature rise in still air (natural convection) is 20°C and in moving air (forced convection) is 8°C. The electrical resistance through the AC power bus is 44mΩ. What chamber temperature should be used for simulating a pedestal still air bottom temperature of 60°C?

1) Calculate the total heat dissipation:

$$\begin{aligned}Q_{UUT} &= Q_{Consumed} + Q_{Ohmic} \\Q_{UUT} &= Q_{Consumed} + I^2 R \\Q_{UUT} &= 70W + (15A)^2(0.044\Omega) \\Q_{UUT} &= 80W\end{aligned}$$

2) Calculate the thermal resistances:

$$\begin{aligned}\Theta_{Natural} &= \frac{\Delta T_{Natural}}{Q_{UUT}} \\ \Theta_{Natural} &= 20^\circ\text{C}/80W = 0.25^\circ\text{C}/W \\ \Theta_{Forced} &= \frac{\Delta T_{Forced}}{Q_{UUT}} \\ \Theta_{Forced} &= 8^\circ\text{C}/80W = 0.10^\circ\text{C}/W\end{aligned}$$

3) Determine the appropriate forced convection chamber temperature:

$$\begin{aligned}T_{Forced} &= Q_{UUT}(\Theta_{Natural} - \Theta_{Forced}) + T_{Natural} \\ T_{Forced} &= 80W(0.25^\circ\text{C}/W - 0.10^\circ\text{C}/W) + 60^\circ\text{C} \\ T_{Forced} &= 12^\circ\text{C} + 60^\circ\text{C} = 72^\circ\text{C}\end{aligned}$$

Application

$\Theta_{Natural}$ and Θ_{Forced} will remain constant for a given housing and chamber combination. Changes to housing geometry or material could affect $\Theta_{Natural}$ and Θ_{Forced} . Natural convection testing in a chamber of different interior dimensions, especially if the clearances between the UUT and the chamber walls are less than 6 inches, will change $\Theta_{Natural}$. Differences in chamber air speed and direction relative to the UUT will change Θ_{Forced} . If the values of $\Theta_{Natural}$ and Θ_{Forced} are in question as the result of a design or chamber change, a test or some other evaluation is recommended.

A simple calculation using Equation 8 will produce corrected chamber temperature for standard forced convection performance

tests. Θ_{Forced} usually varies over a small range for the same UUT tested in different chambers of similar size and air moving capacity. Nodes, amplifiers and line extenders can be tested once in a “pedestal sized” chamber and the resulting values for $\Theta_{Natural}$ and Θ_{Forced} can be applied to tests in similar sized chambers.

Limitations

This test method produces an approximation of pedestal based component operating temperatures in a standard environmental chamber. The test and data reduction procedure described above is a linear approximation of a non-linear system. Internal pedestal temperatures are dependent on many variables including pedestal material, vent design and the local environment. The test procedure has no provisions to simulate solar radiation load. Practitioners must evaluate the heat transfer mechanisms applicable to their situation and then determine the suitability of this approach.

CONCLUSION

Market forces are driving the cable telecommunications industry to more sophisticated network equipment. The heat dissipation and internal temperatures of new equipment will increase if trends in the computer and telecommunications industries are realized. Reduced thermal margins require explicit thermal design requirements and accurate test methods. A test method is presented that allows OSP equipment to be tested in an ordinary environmental chamber in a way that reproduces the component temperatures experienced in a fielded pedestal.

References

- ¹ Branson, K., “More Power! More Power! But where will service providers get it?” Broadband Week, Jan 8, 2001.
- ² Azar, K., “The History of Power Dissipation,” Electronics Cooling Magazine, Jan 2000.
- ³ Steinberg, D., Cooling Techniques for Electronic Equipment, Second Edition, p128.

TRAFFIC ENGINEERING, TRAFFIC CONTROL, PERFORMANCE ANALYSIS AND NODE COMBINING IN DOCSIS-BASED CABLE NETWORKS

Gagan L. Choudhury and Moshe Segal
AT&T Labs, Middletown, NJ 07748, USA
e-mail: {gchoudhury,moshesegal}@att.com
Phone: (732)420-{3721,3713}

Abstract

Simulation and analytic models (based on Markov Chains and Linear/integer Programming) have been developed to analyze and optimally design a cable access system that carries integrated voice/data over IP and DOCSIS. The goal is a design that is cost-effective to network operators and satisfies the QoS needs of each application types.

INTRODUCTION

Over the next few years the United States and the world will see a massive deployment of cable-based access networks to carry high-speed data, voice and video services for consumers and businesses. In some cases the existing cable TV and telephony infrastructure will be upgraded, and in other cases new infrastructure will be deployed. A key to success in the marketplace would be the ability to share common resources among many subscribers (typically few tens to few hundreds) and many applications (many grades of voice, data and video) that satisfies the Quality of Service requirement of each type of subscriber and also becomes cost-effective to the service provider. In this paper we propose traffic engineering, traffic control, performance analysis and node combining methodologies to do the above. We assume that voice and data packets will be transported over the Hybrid-Fiber-Coax (HFC) architecture using IP and DOCSIS (Data over Cable Service Interface Specification) protocols. However, many of the basic observations we make in this paper are likely to be valid even if some other protocol architecture is used.

ARCHITECTURE AND DOCSIS OPERATION

The basic architecture for voice and data is shown in Figure 1. The HFC plant also carries regular television programming (not shown in the figure).

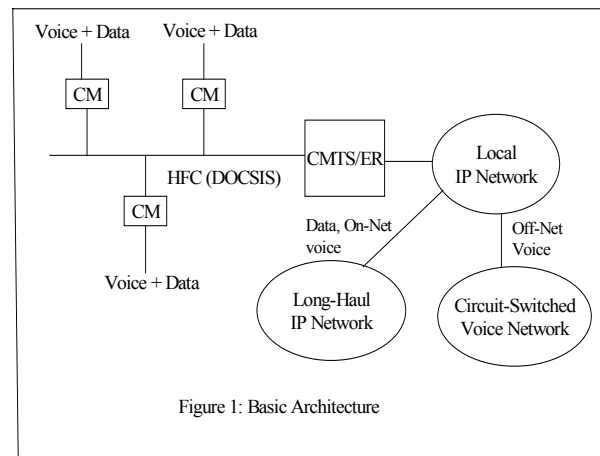


Figure 1: Basic Architecture

Each household has a cable modem (CM) which typically supports several voice lines and one data port. The voice and data are transported as IP packets (typically carrying RTP/UDP payload for voice and TCP or UDP payload for data) from the CM to the CMTS over the HFC infrastructure using the DOCSIS protocol architecture [1,2]. The first version, DOCSIS 1.0, was developed for primarily data applications. The next version, DOCSIS 1.1, has several QoS features for supporting voice and other real-time applications. The CMTS has a built-in edge-router that transports the voice and data traffic to a local IP network. From the local IP network, data typically travels to a long-haul IP network, e.g., the Internet. Voice may stay end-to-end on an IP network (On-

Net voice) or hop-off to a circuit-switched network (Off-Net voice). Usually, the HFC upstream has a frequency range 5-42 MHz (of which typically about 18 MHz is of good quality) and the HFC downstream has a frequency range 54-860 MHz. Both the upstream and the downstream frequency spectra are subdivided into many channels. Typical upstream channel widths are 1.6 MHz and 3.2 MHz even though smaller widths are also possible. Using QPSK modulation, the 1.6 MHz and the 3.2 MHz channels yield data rates of 2.56 Mbps and 5.12 Mbps respectively. Using 16QAM modulation the data rates are twice as much. The downstream channel width is 6 MHz, which yields a data rate of 30.34 Mbps using 64QAM modulation, and a data rate of 42.844 Mbps using 256QAM modulation. Upstream transmission is in bursts separated by guard times and successive bursts may be from different CMs. Each burst has to be a multiple of a basic unit known as a mini-slot that is typically 8 or 16 bytes long. Since upstream transmission is many-to-one, some coordination is needed among the transmitters. This is done by broadcasting allocation Maps periodically over the downstream with each Map precisely defining upstream transmission over a certain period of time in the future. There are three main regions of the Map as shown in Figure 2 below.

Request	Management	Data/Voice/Sig
---------	------------	----------------

Figure 2: Three regions of Map

The request region of the Map allows requests to be sent by all CMs for future transmission opportunities. The contention resolution among requests is similar to what is used in Ethernet [1,2]. The Management region allows management messages to be sent either by individual CMs or by all CMs in contention mode similar to the requests mentioned above. The third region allows grants for data, voice and signaling packet transmission by individual CMs (usually, a CM needs to solicit for such a grant through a successful request packet at an earlier Map). The basic QoS mechanism for tightly controlling delay jitter for voice calls is to provide periodic unsolicited grants to every active voice call. Other QoS mechanisms include real-time polling (periodic unicast request opportunities), non-real-time polling (same as real-

time polling but with lower frequency) and unsolicited grants with activity detection. In order to provide QoS to a particular type of voice or data call it is necessary to define service flow IDs (SFID) on upstream and downstream channels and an associated service ID (SID) on the upstream. A SFID is a particular unidirectional mapping between a CM and the CMTS.

TRAFFIC CHARACTERISTICS

Voice Calls

Voice call arrival process is traditionally modeled as Poisson with a finite or infinite number of sources and that is appropriate in the cable environment as well. Call holding time has a general distribution including exponential, Lognormal or Pareto. The exponential case is easiest to analyze but in many cases the model output (such as blocking probability) is insensitive to the distribution except through the mean and so the other cases may be treated easily as well. In the absence of voice activity detection, voice packets are generated at a constant rate (typically once every 10 or 20 ms) throughout the call duration. In the presence of voice activity detection and silence suppression, there is a sequence of talk-spurts separated by silence periods and packets are generated (at a constant rate) only during the talk-spurts. Here again, silence and talk-spurt distributions may be general but the exponential case is easiest to analyze.

One important way in which the cable environment differs from traditional circuit-switched voice applications is that it is very natural to have multiple data rates. The payload size in a voice packet depends on the packetization interval and encoding scheme, and the overhead in the packet depends on the use or non-use of payload header suppression and the degree of forward error correction done to combat channel noise. The table below shows the data rates over the upstream under some typical scenarios. It is assumed that RTP, UDP, IP, Ethernet, CRC, MAC and physical layer overheads are 12 bytes, 8 bytes, 20 bytes, 14 bytes, 4 bytes, 11 bytes and 24 bytes respectively and under payload header suppression (PHS), all but

two bytes of the UDP, IP and Ethernet layer headers are suppressed. It is also assumed that the mini-slot size is 16 bytes (note that if the payload plus overhead is not an integer multiple of a mini-slot size, some padding is needed to make it so).

Packetization Interval, Encoding scheme	No Payload Header Suppression			With Payload Header Suppression		
	Payload size in Bytes	Overhead in Bytes	Data rate in Kbps	Payload Size in Bytes	Overhead in Bytes	Data rate in Kbps
10 ms, G.711	80	93	140.8	80	53	115.2
10 ms, G.728	20	93	102.4	20	53	64
20 ms, G.711	160	93	102.4	160	53	89.6
20 ms, G.728	40	93	57.6	40	53	38.4

Table 1: Voice Data Rate Variations

The choice of encoding scheme may depend on the voice application (e.g., conversational voice may use G.728 whereas fax and voice-band data may use G.711) or the degree of network congestion (G.711 may be used under light load and G.728 may be used under heavy load). Besides the varying data rates there are other reasons to consider a multi-rate voice model as well. One example is to be able to treat “intra-call” situations where both the called and calling parties use the same channel. In that situation, the “intra-calls” consume twice as much bandwidth but have half as much arrival rate compared to inter-calls whose two end-points use different channels. “Intra-calls” have been analyzed in [10,23].

Data Sessions

It is reasonable to assume that data session arrivals follow a Poisson process (see, e.g., [15,16]). However, the number of packets in a session has an arbitrary distribution that includes heavy-tailed distributions (e.g., Pareto or one of its modified forms). Also, there may be multiple file

transfers within a session with some gap between successive file transfers. The individual packets in a session come at the peak rate that is allowed to each data user (except for the natural flow control in TCP during congestion). Modeling each session as TCP may slow down the simulation and so the impact of TCP is taken into account in an indirect and approximate way. During congestion, file transfer rates are reduced, packets are queued back at the cable modem and DOCSIS serves the various SIDs in roughly a round-robin manner. We assume that TCP instantaneously adapts to this reduced rate. The packet size during the file transfer is limited by the Maximum Transfer Unit (MTU) of Ethernet (about 1500 bytes). We designate these packets as *big*. In addition there are also *small* packets generated due to many reasons such as acknowledgment for downstream traffic, query, and control messages for initiating, maintaining or terminating sessions. The *small* packets arrive according to a Poisson process and they have an arbitrary length distribution bounded upwards by the MTU size and downwards by the smallest allowed packet size.

INTERACTION AMONG VOICE, DATA AND VOICE SIGNALING

Voice is typically given priority over other services over a certain portion of the available channel bandwidth using the unsolicited grant mechanism at an upstream channel and using priority queueing, weighted fair queueing or weighted round robin mechanism at a downstream channel. Therefore, in order to compute voice blocking we typically do not have to consider any other traffic types. In general a voice packet may have to wait for one large data packet transmission. This is typically not a significant issue over the downstream due to its large bandwidth and it may be addressed adequately over the upstream using the fragmentation feature of DOCSIS 1.1. The case of combined DOCSIS 1.0 and 1.1 operation will be addressed later. Also, there are interesting interactions among voice calls that need to be taken into account (e.g., several voice calls may use the same SID).

Delay of data packets and the throughput of data sessions are affected by the presence of voice since data essentially can only use bandwidth left over by voice. Proper engineering of voice traffic is needed to ensure acceptable data performance. Over the upstream, data performance also depends on the way the voice unsolicited grants are allocated since a data packet has to fit in the gap between the unsolicited grants. Since each data packet has to use the request mechanism to get a grant over the upstream, the data performance depends on the efficiency of handling collisions over the request region. The performance of data file transfer may be improved by using the piggybacking mechanism in which a field in the current packet is used as a grant request for the next packet. Finally, the data performance also depends on the maximum rate allowed to a data session and the number of data sessions present simultaneously. The rate available to a data session at a given instant is variable and TCP adjusts itself to the changing rate.

Voice signaling packets contend with data packets over the upstream until the unsolicited grants for the voice stream is established and special mechanisms are needed to ensure acceptable performance of signaling packets in the presence of large data file transfers. Over the downstream, standard QoS mechanisms may be used to ensure acceptable performance of voice signaling packets.

MODELING: SIMULATION VS ANALYSIS

We have developed a simulation model of DOCSIS 1.0 and 1.1 over the upstream that captures the unsolicited grant mechanism for voice transfer and request-based grant mechanism for data transfer. The model can also capture any type of arrival process and call/session duration distribution for voice calls and data sessions. For voice calls the model can also capture the delay/loss performance of adaptive de-jitter buffers. We have also developed several analytic models (faster to run compared to simulations) to capture various aspects of performance. Most of these models are based on Markov chains. Many of the models also have product-form structure resulting in fast computation even in the presence

of many classes of voice and data traffic. In general the models assume Poisson arrival of sessions/calls and exponential session/call duration. However, in many cases the models are insensitive to call/session holding time distribution and the distribution of idle time of a voice/data source and so arbitrary distributions including heavy-tailed ones may be allowed. Besides the Markov-chain based models we have also developed linear and integer programming based models for optimally assigning nodes to CMTS (Cable Modem Termination System) cards. We provide below an example of Markov-chain based modeling. See [23] for other examples. Also, an integrated voice-data model with Poisson arrival of individual voice and data packets has been analyzed in [17-19].

Example Analytic Model: Voice-data Coexistence (Upstream or Downstream)

Let R_T represent the total channel bandwidth available to voice and data, i.e., after subtracting the bandwidth used for management and signaling in the downstream and management, signaling and request contention in the upstream. For the upstream we also assume that the fragmentation feature is used and the fragmentation overhead is accounted for in the data packets. Let $R \leq R_T$ represent the total channel bandwidth available to voice. Let there be q classes of voice calls and s classes of data sessions. Voice call and data session arrival processes are Poisson and the number of packets in a call/session has an arbitrary distribution including a heavy-tailed one [15,16]. Voice calls are blocked if there is not adequate bandwidth, but once admitted, they are guaranteed the data rate r_i (for class i) (using the unsolicited grant mechanism in the upstream). Data sessions can use whatever bandwidth is not being used by voice but individual sessions also have a maximum allowed data rate r_{di} . Data sessions do not have any minimum guaranteed bandwidth and are not blocked. At a given time, if the sum of the data rates for all the voice calls plus the sum of the maximum data rates for all data sessions is less than or equal to R_T then all calls get their requested rates. If however, this requirement is not satisfied then the voice calls continue to get their rates but

one or more data sessions get less than their maximum data rates. To consider the impact of this reduced data rate we consider a *bufferless* model and a *buffered* model. In the *bufferless* model it is assumed that the total duration of the data session is not affected due to this reduced data rate. There are two ways of interpreting this. One way is to assume that the excess data traffic is just lost as in Rate Envelope Multiplexing [20]. The other way is to assume that if the data rate is slower then the user transfers less data and keeps the duration of the data session the same. In the *buffered* model it is assumed that the total number of bytes transferred in a data session is unaffected by congestion and so the data session duration gets elongated.

Bufferless Model

We only consider the finite source version of data sessions and voice calls but the results can easily be extended to the infinite source case by taking appropriate limits. The model has product-form solution (see [3,4]) and insensitivity (except through mean) to both the distribution of session/call duration (which may be heavy-tailed) and the distribution of idle time duration of individual sources.

At first we consider voice calls only since they are unaffected by the presence of data sources. For class i of voice calls we assume that N_i is the number of sources, α_i is the average call arrival rate per idle source, β_i is the offered erlang per source in an infinite-server system, i.e., a system with no blocking, and $a_i = \alpha_i T_i$. Clearly, $\beta_i = a_i / (1 + a_i)$ and $a_i = \beta_i / (1 - \beta_i)$. The Steady-state probability vector $p(\mathbf{n})$ has a product-form solution given by:

$$p(\mathbf{n}) = g(R, \mathbf{N})^{-1} \prod_{i=1}^q \binom{N_i}{n_i} a_i^{n_i} \quad (1)$$

where the normalization constant $g(R, \mathbf{N})$ is given by,

$$g(R, \mathbf{N}) = \sum_{\mathbf{n} \in S(R, \mathbf{N})} \prod_{i=1}^q \binom{N_i}{n_i} a_i^{n_i},$$

$$S(R, \mathbf{N}) = \left\{ 0 \leq n_i \leq N_i, i = 1, \dots, q : \sum_{i=1}^q r_i n_i \leq R \right\} \quad (2)$$

The call blocking probability, P_{bi} for class- i voice calls is given by,

$$P_{bi} = 1 - \frac{g(R - r_i, \mathbf{N} - \mathbf{e}_i)}{g(R, \mathbf{N} - \mathbf{e}_i)} \quad (3)$$

where $\mathbf{e}_i$ is a vector with 1 in the i -th place and 0 everywhere else. The generating function of the normalization constant is given by:

$$G(z, \mathbf{N}) \equiv \sum_{R=0}^{\infty} g(R, \mathbf{N}) z^R = \frac{\prod_{i=1}^q (1 + a_i z^{r_i})^{N_i}}{1 - z} \quad (4)$$

Recursive algorithms for computing $g(R, \mathbf{N})$ may be developed from $G(z, \mathbf{N})$ [12] and as in [5-7,13], highly efficient numerical inversion algorithms may also be developed. Similar algorithms with guaranteed minima and upper limits on the bandwidth usable by different call types have been developed in [11]. Recursive algorithms in the infinite source case have been developed in [8,9].

Next we consider data sessions and assume that class i of data sessions has mean duration T_{di} , N_{di} data sources, each idle source generates a new call at rate α_{di} and $a_{di} \equiv \alpha_{di} T_{di}$. Let $\mathbf{n}_d = (n_{d1}, \dots, n_{ds})$ represent the vector of number of data sessions of various classes and $p_d(\mathbf{n}_d)$ represent its steady-state probability distribution. Since there is no blocking of data sessions and the data session duration is unaffected by congestion, $p_d(\mathbf{n}_d)$ is given by the expression

$$p_d(\mathbf{n}_d) = \prod_{j=1}^s (1 + a_{dj})^{-N_{dj}} \binom{N_{dj}}{n_{dj}} a_{dj}^{n_{dj}} \quad (5)$$

We define the probability of no data degradation, P_{ndd} , as the probability that all the data calls get their respective maximum rates and

$P_{dd} \equiv 1 - P_{n\ddot{d}d}$ as the probability of data degradation. In order to engineer a channel for combined voice-data service we use the voice call blocking probability as the performance measure for voice and the probability of data degradation as the performance measure for data. (Other performance measures such as the mean rate for data calls of class i may also be derived).

$$P_{n\ddot{d}d} = 1 - P_{dd} = \text{Pr ob.} \left(\sum_{i=1}^q r_i n_i + \sum_{j=1}^s r_{dj} n_{dj} \leq R_T \right)$$

$$= \sum_{\substack{\mathbf{n} \in S(R, \mathbf{N}) \\ (\mathbf{n}, \mathbf{n}_d) \in S_T(R_T, \mathbf{N}, \mathbf{N}_d)}} p(\mathbf{n}) \cdot p_d(\mathbf{n}_d) \quad (6)$$

Where, $S(R, \mathbf{N})$ is as in (2) and

$$S_T(R_T, \mathbf{N}, \mathbf{N}_d) = \left\{ \begin{array}{l} 0 \leq n_i \leq N_i, i=1, \dots, q; 0 \leq n_{dj} \leq N_{dj}, \\ j=1, \dots, s; \sum_{i=1}^q r_i n_i + \sum_{j=1}^s r_{dj} n_{dj} \leq R_T \end{array} \right\} \quad (7)$$

For small q, s, R , and R_T it is possible to compute $P_{n\ddot{d}d}$ directly from (6) but otherwise we need to develop more efficient algorithms. Using (6), (5), (1) and (2) it can be seen that

$$P_{n\ddot{d}d} = \prod_{j=1}^s (1 + a_{dj})^{-N_{dj}} g_T(R, R_T, \mathbf{N}, \mathbf{N}_d) / g(R, \mathbf{N}) \quad (8)$$

where $g(R, \mathbf{N})$ is as in (2) and the second normalization constant is given by

$$g_T(R, R_T, \mathbf{N}, \mathbf{N}_d) = \sum_{\substack{\mathbf{n} \in S(R, \mathbf{N}) \\ (\mathbf{n}, \mathbf{n}_d) \in S_T(R_T, \mathbf{N}, \mathbf{N}_d)}} \prod_{i=1}^q \binom{N_i}{n_i} a_i^{n_i} \prod_{j=1}^s \binom{N_{dj}}{n_{dj}} a_{dj}^{n_{dj}} \quad (9)$$

Next we obtain the generating function as

$$G_T(z_1, z_2, \mathbf{N}, \mathbf{N}_d) \equiv \sum_{R=0}^{\infty} \sum_{R_T=0}^{\infty} g_T(R, R_T, \mathbf{N}, \mathbf{N}_d) z_1^R z_2^{R_T}$$

$$= \frac{\prod_{i=1}^q (1 + a_i(z_1 z_2)^{r_i})^{N_i} \prod_{j=1}^s (1 + a_{dj} z_2^{r_{dj}})^{N_{dj}}}{(1 - z_1)(1 - z_2)} \quad (10)$$

We obtain $g_T(R, R_T, \mathbf{N}, \mathbf{N}_d)$ by numerically inverting Equation (10) using the methodology in

[5-7], and then obtain $P_{n\ddot{d}d}$ from (8) even for large q, s, R and R_T .

Buffered Model

This model will be explored in detail in a future paper and here we just consider the simple case of one class of voice and one class of data calls. Furthermore, we assume each class to have infinite number of sources, Poisson arrivals and exponential holding times. In general this model does not have a product-form structure and is sensitive to the entire holding time distribution (even though our simulation study, not reported here, shows that the sensitivity is often not strong). Let λ and λ_d represent the arrival rates for voice and data calls, T and T_d represent their mean holding times, and $\mu = 1/T$, $\mu_d = 1/T_d$. Also, let r and r_d represent the data rates for voice and data calls. Let n and n_d represent the number of voice and data calls at a given instant. The state space is given by

$$S(R) = (0 \leq n \leq \left\lfloor \frac{R}{r} \right\rfloor, 0 \leq n_d < \infty) \quad (11)$$

The state vector (n, n_d) follows a continuous time Markov chain with transition rates as given below:

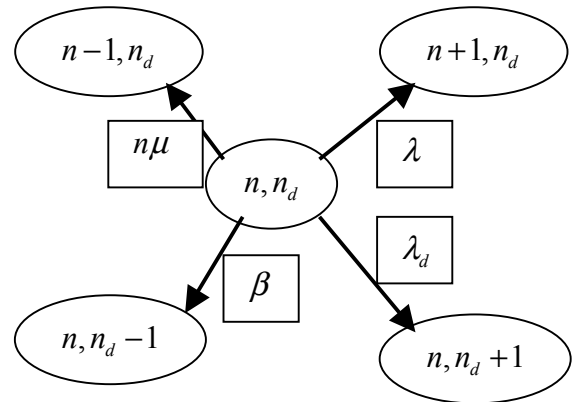


Figure 3: Markov Chain Transition Rates For the Buffered Model

Where,

$$\beta = \min \left(n_d \mu_d, \frac{R - nr}{r_d} \cdot \mu_d \right) \quad (12)$$

Note that the second term within the minimum expression in (12) accounts for the rate reduction and associated elongation of data call holding time under congestion and is not there in the *bufferless* model. One difficulty in solving the Markov chain above is the infinite upper bound on n_d . In

practice however an upper bound n_d^u can be placed on n_d above which state probabilities are too small or the data rate to data calls is so small that new data calls are refused. Once we solve the Markov chain, we can evaluate various performance measures such as the mean rate for data calls or the data degradation probability defined earlier.

NUMERICAL EXAMPLES

Some Generic Assumptions

- Upstream Channel:
 - 1.6/3.2 MHz RF, QPSK modulation, 2.56/5.12 Mbps.
 - Minislot size = 8 Bytes.
 - Contention and management regions each use 4% bandwidth on the average.
 - One request packet fits in a minislot
 - Map interval: Between 2 and 10 ms.
- Voice calls
 - 10 ms packetization.
 - G.711 with PHS.
 - Payload + overhead = 136 Bytes.
 - Poisson arrivals, exponential call holding time with mean 3 minutes.
 - 2.5 lines per subscriber, 0.14 Erlangs per line
- Data/Signaling packet overheads
 - All layer 3 and above overheads (IP and TCP or UDP) included in the packet size.
 - Layer 2 (Ethernet + MAC) overhead: 24 bytes.

- If fragmented then fragmentation overhead: 16 Bytes.
- Layer 1 (Physical) overhead: 34 Bytes.
- Data Sessions
 - Two types of packets
 - *Big*: 1500 bytes (Max allowed, File Transfers).
 - *Small* : 94 Bytes average, Coefficient of Variation 1 (Acknowledgments, Query, Control, Interactive, etc.)
 - 80% of the Packets are *Small*.
 - 80% of the Bytes Come from *Large* Packets.
 - *Big* packets come from file transfers with mean file size of 15 KBytes, Coefficient of Variation 2
 - Packet size: 64 Bytes (including Layer 3 and above overhead).
- Simulation Runs (in case of Simulation):
 - Each run captures 3 hours of operation with the first hour used for warm-up and statistics is collected only over the last two hours.
 - Each run is repeated 11 times. The first run is used to get a rough estimate of the 95th and 99th percentiles and the last 10 are used for actual computations including confidence intervals (not shown).

Voice-Data Integration: Hard Vs. Soft Partitioning, Impact of Packing of Voice Unsolicited Grants

We consider the following partitioning schemes among voice and data and packing schemes for voice unsolicited grants:

- Hard Partition: Voice and data/signaling use separate bandwidth regions with no sharing.
- Soft Partition: Data/signaling can use currently unused voice bandwidth. Two cases are considered
 - Random packing of voice unsolicited grants within each Map interval.
 - Pack unsolicited grants away from the data region and close holes within each Map interval. Closing of holes is done by potentially moving voice unsolicited grants

of existing calls as a call leaves (a scheme without closing holes has been analyzed in [14]).

We assume that fragmentation feature of DOCSIS 1.1 is used, upstream channel bandwidth is 1.6 MHz, at most 18 simultaneous voice calls (this uses up about 83% of bandwidth available to voice and data), 74 voice lines and data traffic being generated from 10 cable modems. Figure 4 below shows the performance of the three schemes. It is clear that the performance is substantially better for the soft partition schemes and among the two soft partition schemes, performance is better if all the voice unsolicited grants are packed at one side as one contiguous block thereby leaving one contiguous block for data and signaling.

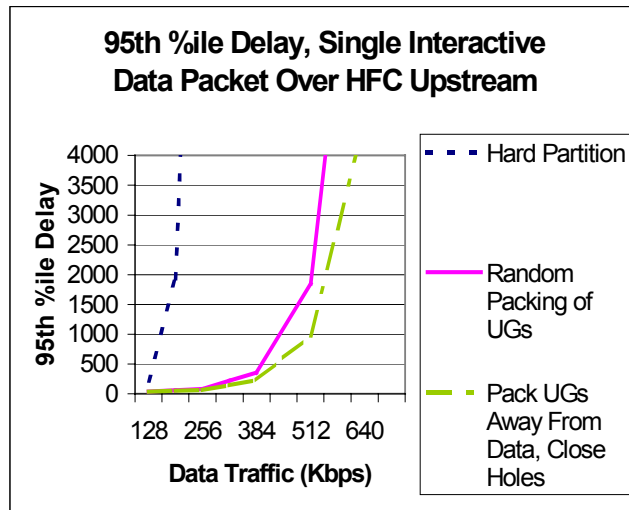


Figure 4

DOCSIS 1.0 Modems for Data and DOCSIS 1.1 Modems for Voice

DOCSIS 1.1 modems do not support fragmentation and so either the entire data packet or nothing can be transmitted in the gap left over by voice unsolicited grants. We at first assume that no jittering of voice packets is allowed to make some extra room for data packet transmission. Figures 5 and 6 show the performance of the two voice packing schemes under this condition and with MTU sizes of 500 and 1500 bytes respectively. It is assumed that upstream channel bandwidth is 3.2 MHz and there are 135 data subscribers each being active for about 4.5% of the time. It is seen that the difference between the two packing schemes is

much more significant compared to the pure DOCSIS 1.1 case in Figure 4.

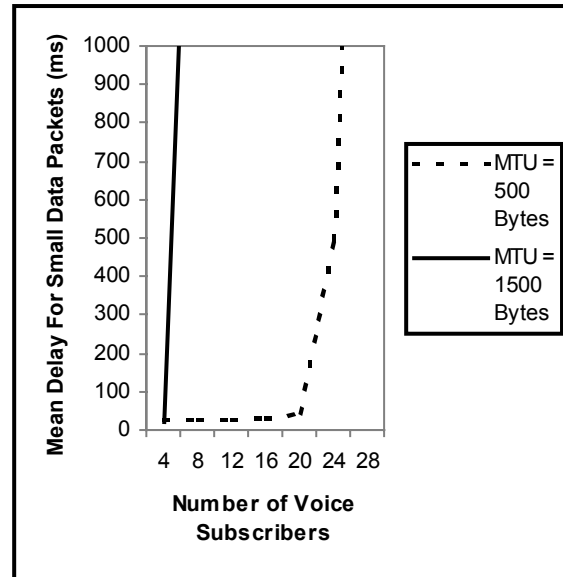


Figure 5: No Fragmentation Used, Pack Random

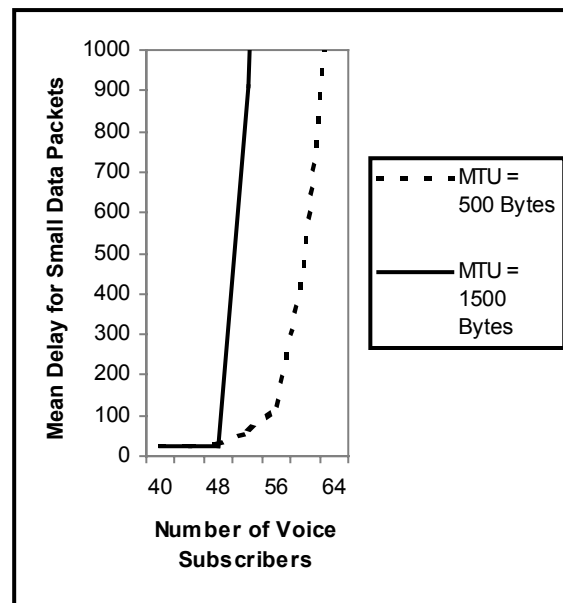


Figure 6: No Fragmentation Used, Pack UGs Away From Data and Close Holes

The performance of the “Pack UGs Away From Data and Close Holes” scheme can be further improved by allowing some delay jitters to voice calls. This is illustrated in Figure 7.

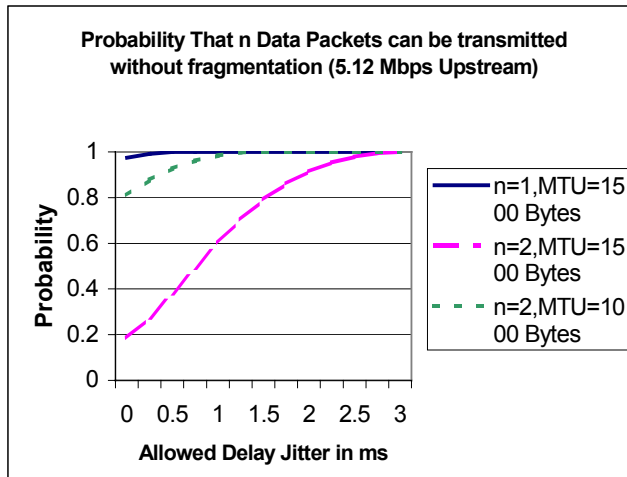


Figure 7

Voice Signaling Performance in the Presence of Data and Impact of Map Frequency

As mentioned earlier, the voice signaling packets ride with data over the HFC upstream until the unsolicited grants for the voice RTP stream is established. Figures 8-10 below consider three cases:

- Case 1: Data and signaling use the same SID.
- Case 2: Data and signaling use different SIDs but once requests of either kind arrive successfully at CMTS, they are treated at the same priority.
- Case 3: CMTS serves signaling at a higher priority (in addition to having separate SIDs for signaling and data).

Assumptions on voice and data are the same as in the first numerical example of voice-data integration and it is further assumed that signaling packets are 100 bytes long (including overhead) and only 2% of the (data + signaling) traffic is signaling (in terms of bytes). It is seen that signaling delay can be in the hundreds of ms to even seconds (depending on data traffic volume) in Case 1. This is quite unacceptable since the signaling delay may be encountered several times in the overall critical path of signaling. In case 2 the delay gets down to below 100 ms and in Case 3 the delay is only a few tens of ms even in the presence of heavy data traffic. The data traffic performance (not shown) is practically the same in all three cases due to the light volume of signaling. Therefore, it is preferable to use a separate SID for

signaling and also give it higher priority compared to data at the CMTS.

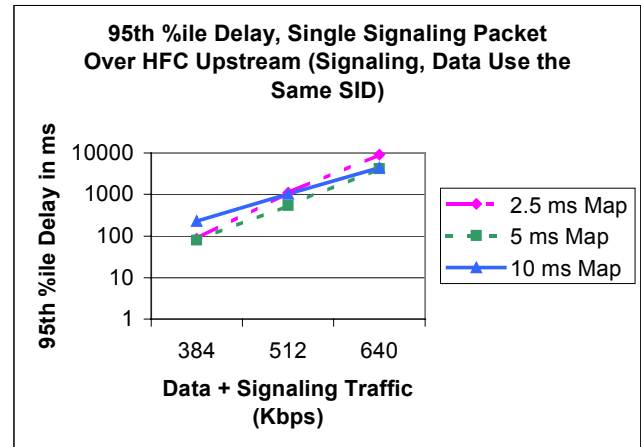


Figure 8

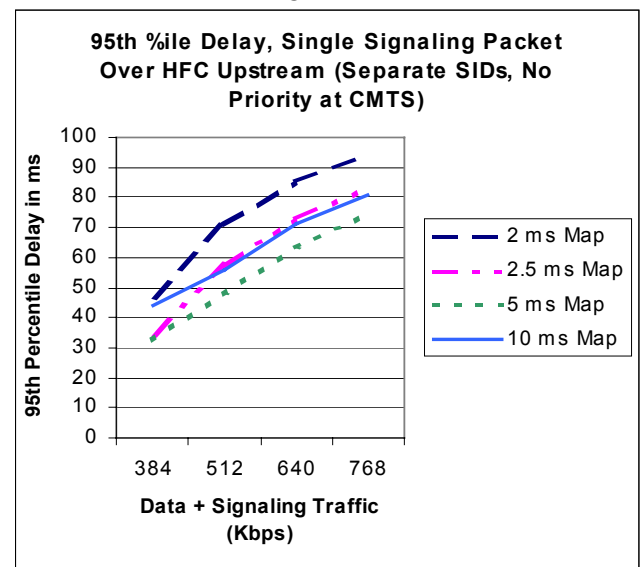


Figure 9

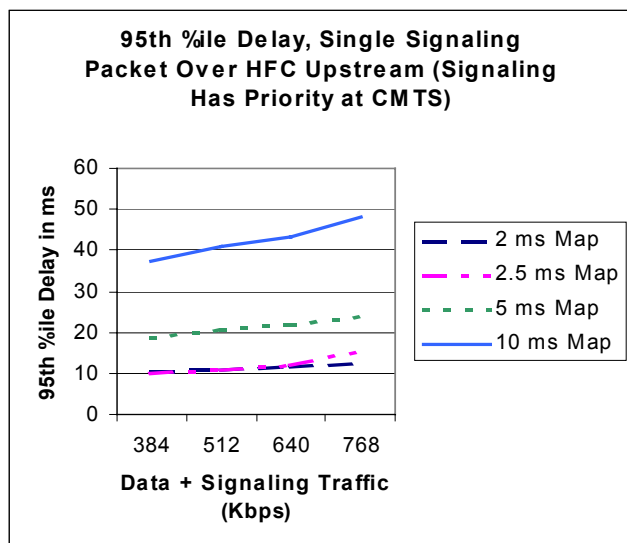


Figure 10

Voice-Path Latency Vs HFC Upstream Capacity

The voice path latency may be reduced by reducing the packetization interval but such a reduction has an adverse impact on the HFC upstream capacity. This tradeoff is shown in Tables 2 and 3. For voice-path latency computation it is assumed that with a 10 ms packetization, the typical/high estimates of round-trip delay at the cable access (or egress) is 40/55 ms for an end-to-end IP call and 70/90 ms if the call has to go over to the PSTN. Furthermore it is assumed that typical/high estimates of round-trip delays at PSTN access, PSTN backbone and IP backbone for a coast-to-coast call are 6/10 ms, 68/94 ms and 95/115 ms respectively. For HFC upstream capacity computation it is assumed that channel RF bandwidth is 1.6 MHz, 25% of bandwidth set aside for data/request/management/signaling, 0.5% call blocking probability, 30% take-rate for voice, efficiency of upstream channel pooling 10% less compared to full access and G.711 with PHS. Based on Tables 2 and 3 it appears that a reasonable compromise between low voice-path latency and high HFC upstream capacity would be to assume a packetization interval of 10 to 15 ms.

Type of Call	Packetization Interval			
	5 ms	10 ms	15 ms	20 ms
End-to-End IP	165/215	175/225	185/235	195/245
PSTN Backbone/Egress, Cable Access	134/184	144/194	154/204	164/214
PSTN Backbone, Cable Access/Egress	188/254	208/274	228/294	248/314

Table 2: Typical/High Round-trip voice-path latency for Coast-to-coast calls (in ms)

Packetization Interval	# of Simultaneous Calls Per Upstream	# of Households Passed (6 Upstreams, no Channel Pooling)	# of Households Passed (6 Upstreams and Channel Pooling)
5 ms	12	300	480
10 ms	17	500	716
15 ms	20	640	860
20 ms	22	720	956

Table 3: HFC Upstream Capacity

Impact of HFC Upstream Capacity on PHS, Compression, Modulation and Channel Width

The impacts are shown in Table 4. The assumptions are the same as in Table 3 except that no upstream channel pooling is considered and it is assumed that total available upstream RF bandwidth is 9.6 MHz, which can be partitioned into 1.6 MHz or 3.2 MHz RF channels. Also, both QPSK and 16QAM modulations are considered.

Encoding Scheme and Header Suppression	Capacity in # of Households Passed			
	Channel Width 1.6 MHz		Channel Width 3.2 MHz	
	QPSK	16QAM	QPSK	16QAM
G.711, No PHS	340	860	470	1030
G.711, PHS	500	1140	660	1360
G.728, No PHS	680	1420	830	1670
G.728, PHS	1240	2260	1440	2570

Table 4

A Methodology for Combining Nodes to a CMTS

Typically, a CMTS or each card of a CMTS supports a certain number of downstream and upstream channels. Using the methodology in Section 5 we can compute the number of voice and data subscribers that can be supported in each such channel and based on market penetration assumptions those numbers can be converted to the number of HHPs (households passed) that can be supported per CMTS card. Suppose in a given neighborhood there are N nodes and the i -th node has $H(i)$ households passed. The problem we address in this section is how to find an assignment of nodes to CMTS cards that satisfies the following:

- A CMTS card is assigned no more than SMAX (e.g., 4) nodes,
- A CMTS card is assigned no more than HMAX (e.g. 2500) HHPs

Optionally, for any feasible assignment of nodes to CMTS cards and for every card, calculate left-over HHP margins and the smallest one among them and then find an assignment that maximizes the smallest margin value. The rationale for doing this is to keep the maximal amount of idleness at each CMTS card for future growth. This is a “packing” type problem. We first need to estimate the number of CMTS cards, C , by rounding up

$\sum_{i=1}^N H(i) / \text{HMAX}$ to the closest integer. The

distribution of the $H(i)$ ’s and their Size compared to HMAX could result in breakage and require increasing the value of C until a feasible solution is found. In what follows we formulate a Linear Binary Program to solve this problem.

Variables:

- $X(i,c)$: assignment variable. Equals 1 when node ‘i’ is assigned to card ‘c’ and equals 0 otherwise.
- $y(c) \geq 0$: left over HHP margin in card ‘c’
- w : A weight factor which assumes the value 1 for optimal solution and 0 for feasible solution

Constraints:

- $\sum_{c=1}^C x(i,c) = 1$ for all ‘i’ every node must be assigned to some card
- $\sum_{i=1}^N x(i,c) \leq \text{SMAX}$ for all ‘c’ a card serves at most SMAX nodes
- $\sum_{i=1}^N x(i,c) \cdot H(i) + y(c) = \text{HMAX}$ for all ‘c’ a card serves at most HMAX HHP
- $y(c) \geq z$ for all ‘c’ all cards have at least a margin of z HHP
- $z \geq \text{YMAX}$ the minimum margin is at least YMAX

Objective:

- Maximize $w \cdot z$

We implemented the model with AMPL [21] and CPLEX [22]. For small problems (e.g., $C = 10$) we can obtain both the feasible and optimal solutions very quickly by setting $w = 1$ and $\text{YMAX} = 0$. For large problems (e.g., $C=70$) we can still find the feasible solution quickly by setting $w = 0$ and an estimate for YMAX but the optimal solution takes too long to verify. With consistent unit changes, subscribers can be used in place of HHPs in this model. For illustration purpose, we present in Figure 11 an optimal solution to a 10-node 5-CMTS-cards problem.

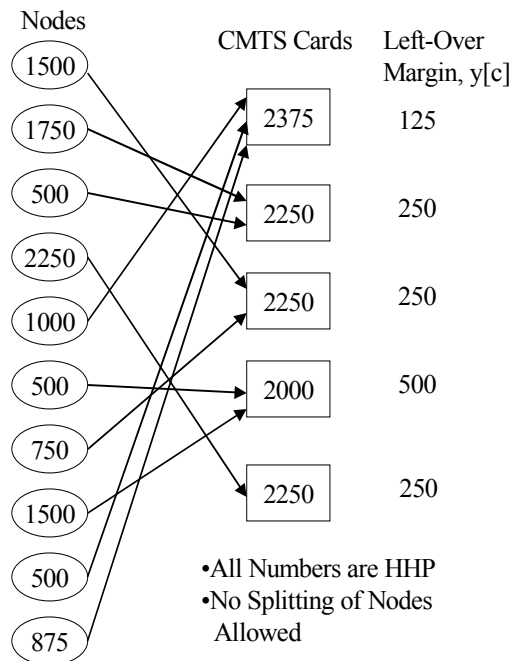


Figure 11. Optimal Assignment of 10 Nodes to 5 CMTS Cards

CONCLUSIONS

We have developed analytic and simulation models for the purpose of engineering, analyzing and controlling combined voice and data traffic over the HFC architecture. We show that the overall system performance and capacity are strongly affected by the modulation and encoding schemes, DOCSIS features such as fragmentation and payload header suppression, and the sharing technique between voice and data.

REFERENCES

- [1] *Data-Over-Cable Service Interface Specifications, Radio Frequency Interface Specifications 1.0*, SP-RFI-I05-991105, November 1999, Cable Television Laboratories, Inc.
- [2] *Data-Over-Cable Service Interface Specifications, Radio Frequency Interface Specifications 1.1*, SP-RFIv1.1-I05-00714, July 2000, Cable Television Laboratories, Inc.
- [3] E. Brockmeyer, H.L. Halstrom and Arns Jensen, "The life and works of A.K. Erlang", The

Copenhagen Telephone Company, Copenhagen 1948.

- [4] F. P. Kelly, *Reversibility and Stochastic Networks*, Wiley, New York, 1979

- [5] G. L. Choudhury, D. Lucantoni and W. Whitt, "Multi-Dimensional Transform Inversion with applications to the transient M/G/1 Queue," *Annals of Applied Probability*, vol. 4, pp. 719-740, 1994

- [6] G. L. Choudhury, K. K. Leung and W. Whitt, "An Algorithm to Compute Blocking Probabilities in Multi-rate Multi-Class Multi-Resource Loss Models," *Advances in Applied Probability* 27, , pp. 1104-1143, December, 1995.

- [7] G. L. Choudhury, K. K. Leung and W. Whitt, "An Inversion Algorithm to Compute Blocking Probabilities in Loss Networks with State-Dependent Rates," *IEEE Transactions on Networking*, Vol. 3, No. 5, pp. 585-601, October 1995.

- [8] J. S. Kaufman, "Blocking in a Shared Resource Environment," *IEEE Transactions on Communications*, 29, pp. 1474-1481. 1981.

- [9] J. W. Roberts, "A Service System With Heterogeneous User Requirements," In *Performance of Data Commun. Systems and their Applications*, ed. G. Pujolle, pp. 423-431, North Holland, Amsterdam, 1981.

- [10] A. W. Jones, "Measurement of Non-Random Telephone Traffic with Special Reference to Line Concentrators," *Proceedings of the 3rd International Teletraffic Congress*, Paris, France, September 11-16, 1961.

- [11] G. L. Choudhury, K.K. Leung and W. Whitt, "Efficiently Providing Multiple Grades of Service with Protection Against Overloads in Shared Resources," *AT&T Technical Journal*, Volume 74, No. 4, pp. 50-63, July/August, 1995.

- [12] L.W.N. Delbrouck, "On the Steady-state distribution in a service facility carrying mixtures of traffic with different peakedness factors and capacity requirements," *IEEE Trans. On Comm.*, 31, pp. 1209-1211, 1983.

- [13] G. L. Choudhury and W. Whitt, "Q-Squared: a New Performance Analysis Tool Exploiting Numerical Transform Inversion," *Demonstration at the 15th ITC*, Washington, D.C., USA, June 23-27, 1997.

- [14] W-S. Lai, "DOCSIS-Based Cable Networks: Impact of Data Packets on Upstream Capacity," 14th ITC Specialists Seminar on Access

Networks and Systems, Barcelona, Spain, April 25-27, 2001.

[15] V. Paxson and S. Floyd, "Wide Area Traffic: The failure of Poisson Modeling," IEEE/ACM Trans. On Networking 3, pp. 226-244, 1995.

[16] M. Crovella and A. Bestavros, "Self-Similarity in World Wide Web Traffic: Evidence and Possible Cause," in: *ACM Sigmetrics '96*, 1996.

[17] M. Segal, "A preemptive priority model with two classes of customers, Proc. ACM/IEEE Second Symposium on Problems in the Optimization of Data Communications Systems," Palo Alto, Calif., October 20-22, 1971, IEEE Cat. No. 71C59-C, pp. 168-174.

[18] S. Halfin and M. Segal, "A priority queueing model for a mixture of two types of customers," SIAM J. Appl. Math., Vol. 23, No. 3 November 1972, pp. 369-379.

[19] S. Halfin, "Steady-state distribution for the buffer content of an M/G/1 queue with varying service rate," SIAM J. Appl. Math., Vol. 23, No. 3 November 1972, pp. 356-363.

[20]

[21] J. W. Roberts (Ed.), "Performance Evaluation and Design of Multiservice Networks," Final report of COST 224 project, Office for official publications of the European Communities, Luxembourg, 1991.

[22] R. Fourier, D. M. Gay and B. Kernighan, "AMPL, A Modeling Language for Mathematical Programming," The Scientific Press, 1993.

[23] *Using the CPLEX Callable Library, Version 4.0*, 1995.

[24] G. L. Choudhury and M. Segal, "Cable Infrastructure to Carry Voice and Data – Performance Analysis, Traffic Engineering and Cable Head-end Node Design", 14th ITC Specialists Seminar on Access Networks and Systems, Barcelona, Spain, April 25-27, 2001.

2. Acknowledgements

We would like to thank Arun Guha, Wai Sum Lai, Bill Lester and Doug Nortz for their valuable input and comments.

USE OF MPEG-4 IN BROADBAND APPLICATIONS

Mukta L. Kar, Ph.D. and Yasser Syed, Ph.D.
Cable Television Laboratories, Inc.
Sam Narasimhan, Ph.D. and Ajay Luthra, Ph.D.
Motorola Broadband Communication Sector

ABSTRACT

The emerging MPEG-4 standard in development today will provide new opportunities to broadband environments. Streaming Media, enabled through MPEG-4 technologies, can create personal viewing and interactive applications, such as Video-on-Demand (VoD), targeted advertising, and hot-button interactive product purchasing. This paper will describe some of the newer capabilities added to the MPEG-4 standard that are advantageous to broadband environments. Furthermore, cable initiatives for digital local ad insertion (DVS 253 and DVS 380) will be discussed in light of digital compression technologies, such as rate-remultiplexing. We will look at how MPEG-4 can achieve this within an MPEG-2 structure and still allow for bandwidth reduction. Finally, this paper will discuss how a targeted advertising capability can result from MPEG-4 implementation.

Introduction

Since the early days of cable, the primary source of revenue has been the broadcast delivery of premium analog content. This type of delivery mechanism does not utilize resources and information in an optimized way. Resources and information are wasted in this approach because not everyone will watch programs in their entirety at a pre-scheduled time or consider the information (advertisements) as relevant. Many cable systems have been upgraded from one-way analog to the modern hybrid-fiber coax (HFC) two-way digital broadband networks of today. Digital delivery using advanced digital modulation and compression technology has expanded a limited analog RF (Radio Frequency) spectrum to a large digital bandwidth potentially increasing the channel capacity of a cable plant by a few folds [6]. This digital bandwidth can be used to offer more niche channels to capture a more targeted audience interested in the programs as well as the advertisements being

shown. This is one of the few ways to utilize digital delivery optimally to increase the revenue of the cable industry by offering more niche programming/content. Additionally, the large digital bandwidth also can be used for delivering Internet protocol (IP)-based services [5][6].

Streaming Media adds a new dimension by allowing a personally designed viewing experience. The viewer's choice is not limited to local headend offerings. With streaming media, a viewer can request content at a time when he/she desires to see it, and from a source that may reside anywhere in the world. This content can be delivered through a headend-managed, DOCSIS-enabled IP transport stream or through a traditional MPEG-2 delivery system to a set-top box (STB) [6]. It also adds in the ability to target content to a specific type of viewer rather than broadcasting the same content to the entire audience. A streaming media session adds increased value by capturing a viewer who is definitely interested in the program. Revenue can come from the customer actually purchasing the streaming media content or through the use of targeted advertisements that are more relevant to the viewer [5].

Through Streaming Media, the customer can actually purchase content that he/she is willing to watch at his/her convenience. One of the most visible applications of this technology is video-on-demand (VoD) services where movies can be streamed to the customer at his request. The application can be extended to other types of content, such as "how to" programs, news programs, and weekly serials. This service can be enhanced by adding in VCR capabilities, such as pause, fast-forward, and rewind. However, there would be an increased amount of downstream data generated by each customer request [4], and some upstream traffic due to interactive VCR features. The additional traffic in both directions would be an added strain to a

two-way cable plant. For services like these to exist on a large scale in a two-way cable infrastructure, a technology like MPEG-4 needs to be implemented to permit high-quality interactive video to stream at low bit rates [4].

Another revenue source opened up by streaming media is targeted advertising. In the existing method of broadcast delivery, national commercials are inserted into the program stream by the networks before delivery to the cable headends and broadcast affiliates. However, advertisements and short programs originated locally are being inserted (overwritten on top of selected national ads) per cue tones delivered by the networks. The revenue from the commercials is not only significant compared to the total revenue, but provides a good cash flow to the cable industry. The existing method of ad delivery is not very efficient. Technically, the existing approaches splice the local ad into the national distributed program stream using analog techniques by uncompressing and converting the digital streams into an analog format. The output is then delivered over the local cable network as an analog signal with a degradation in visual quality as well as a loss in bandwidth revenue. In terms of scope, these commercials are broadcast to the entire regional viewing audience (large area), but fail to be relevant to many in the viewing audience. With streaming combined with digital statistical remultiplexing techniques, these slots can contain different local digitally inserted commercials that can be targeted towards each customer [5]. For instance, a car company can put out a different car commercial for each of their product lines based upon a customer profile for the same commercial time slot. This way the most relevant commercial reaches the customer. As techniques become further developed, streaming also can enhance broadcast commercials by combining additional data to create a higher visual quality commercial instead of having the commercial constrained to the encoding bit rate of the channel. Lastly, the commercial can be made more interactive by providing a way for the customer to stream back his responses. An example of this would be the ability to immediately buy a product advertised in a commercial. Similar to VoD, these types of services will increase traffic loads in the cable network to accommodate the individualized information generated by the streaming application. A technology like MPEG-4 will be necessary to allow these services to exist on a large scale by reducing the bandwidth required

for low bit rate transmissions and allowing for interactivity [1][4].

Streaming Media in MPEG-4

MPEG-4 Background

The MPEG-1 standard was created to satisfy the need for storage applications, such as CD-ROMs and adopted frame-based video compression methods. Later MPEG-2 standards were developed to meet requirements of the broadcast industry. It included field/frame-based video coding to deal with interlaced video as well. In brief, the MPEG-2 standard was optimized primarily for one-way broadcast delivery of television content [1]. The MPEG-4 standard effort initially differed by focusing on low bitrate coding applications over IP connections. Later its scope increased to cover a wider range of multimedia applications including videophone, interactive TV, streaming video etc. [4]. Additionally, to include interactive capability at various levels, the coding paradigm changed from frame-based to content-based or object-based coding. One of the more important advantages of MPEG-4 is that it is well optimized for scalable low bit rate (LBR) transmissions (less than 2 Mbps) and the second one is the ability to selectively incorporate natural and synthetic objects into a scene. These advantages will allow streaming media to exist on a large scale by economizing the available bandwidth while allowing interactivity to be integrated [3].

Object-based coding has other advantages. It reduces bandwidth demands by allowing for object-based coding instead of frame-based coding that exists in MPEG-2 technologies. In this case, instead of a frame being coded in its entirety, separate audio/video (A/V) objects within the frame can be encoded at different quality and rates. Objects do not have to be re-transmitted each time a scene changes, rather only manipulation information (scale, translation) of the object could be sent. Each A/V element can be encoded in its own elementary or set of elementary streams in the form of video object planes (VOP). Different encoding parameters can be assigned based upon the nature of the object to allow for the most efficient type of encoding. For instance, a natural image can be encoded by a video codec while a text-based object can be encoded by a text-based codec, which would display a higher quality object

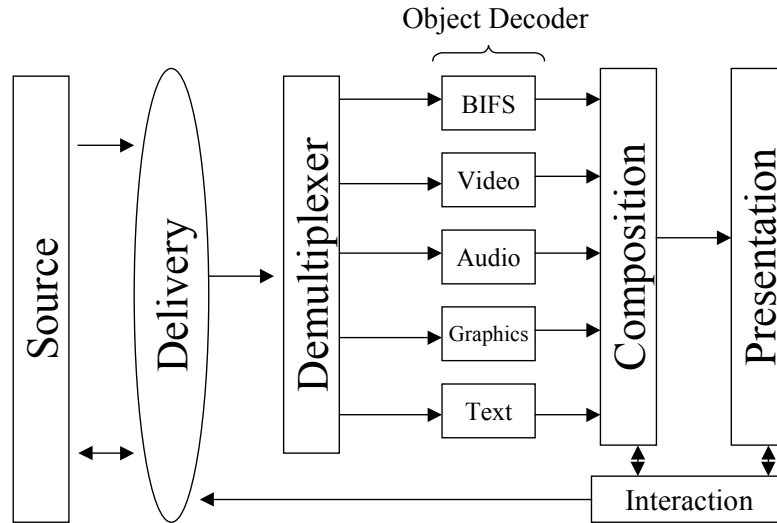


Figure 1. Block Diagram of MPEG-4 System

with a smaller amount of bits. By using the right codec type, bits can be conserved without sacrificing video quality. There are five general categories of coders (Figure 1): 1) video, 2) audio, 3) graphics, 4) text, and 5) scene. MPEG-4 adapts each of its codecs to conform to multiple profiles and levels of transmission to accommodate different delivery formats. The video and audio codecs are used on natural video and audio objects optimized for good quality at low bit rates. In addition, MPEG-4 adds in a spatial and temporal scalability factor to traditional A/V coding and provides graceful degradation of objects during times of congestion [3]. The graphics coder provides a means to animate and render synthetic objects. Synthetic objects can be computer-generated objects with interactive components. An example of this would be a “hot-button” to purchase an advertised product. The text coder provides an efficient way to code text. The scene coder, called BIFS (Binary Format for Scene), is responsible for scene composition and rendering. This coder can manipulate (spatially, temporally), layer and even edit out objects from the scene [7]. It also can add or delete streams such that a directed channel change (DCC) can occur to allow for targeted advertising [12].

New Scalable Profiles in MPEG-4

The adoption of the advanced simple profile (ASP) and fine grain scalability (FGS) profile in the MPEG-4 visual standard allows for different layered levels of quality that are advantageous for streaming media over the Internet, which consists of many heterogeneous networks. ASP allows the highest possible quality within the MPEG-4 standard for a traditional video consisting of rectangular shape frames by allowing the use of B-frames, quarter-pel interpolation, global motion compensation (GMC) and interlaced coding format. It also allows pictures to be compressed at higher than common interchange format (CIF) resolution—half horizontal resolution (HHR) and full resolution—and at high compressed bit rates. It does not include shape coding tools and thus does not have the complexity associated with arbitrary shape coding [1][3]. The addition of FGS provides for transmission of base and enhancement layer streams. The base layer contains the lowest quality coded image of the object and is compliant with the ASP. The enhancement layer in the FGS profile adds to the base layer to increase the visual quality of the object. A spatial enhancement layer can be done by the FGS layer, and a temporal enhancement layer to the visual stream can be

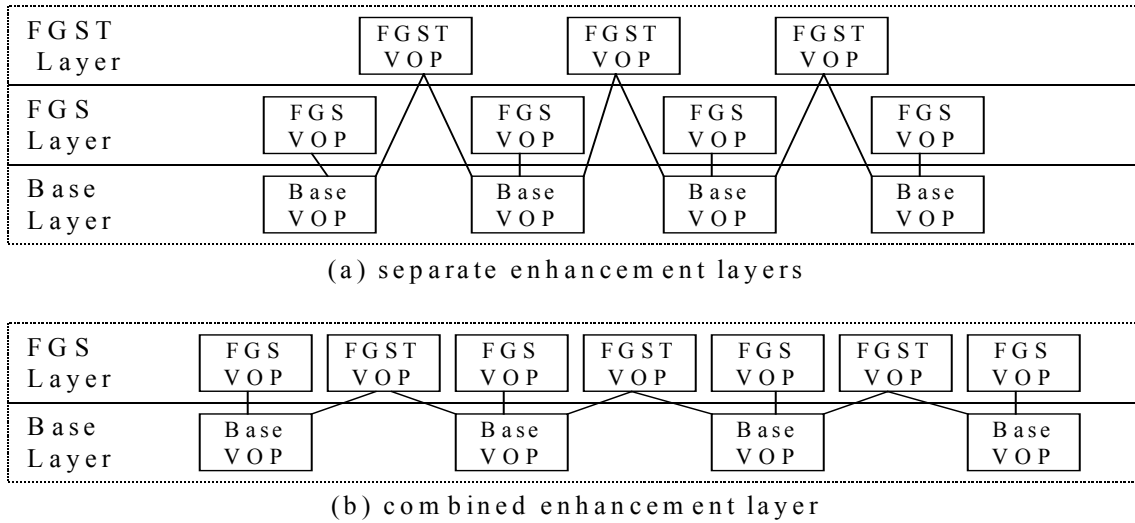


Figure 2. Base & Enhancement Layers Combinations for Streaming Video over the Internet

done by the FGS temporal scalability (FGTS) layer. Both these layers only require the base layer information to be decoded and have a scalable embedded bitstream property (where the number of bits processed is proportional to the image quality). This is different from the discrete scalability profile put forth in MPEG-2 or earlier versions of the MPEG-4 profiles. Each type of enhancement layer can be layered on top of each other or combined into a single enhancement stream layer (Figure 2) The different profile levels (five for the ASP and five for FGS) adds in different capabilities, such as number of objects, visual image size, temporal and/or spatial scalability, buffer size, maximum packet length, and maximum bit rate [2].

This base/enhancement layer partition provides several advantages. In an IP-congested environment, a graceful degradation of the image (by sacrificing the quality of certain visual objects) can occur during times of net congestion through the use of an embedded bit stream property. Also, multiple CPE devices with different bandwidths, QoS parameters, or processing restrictions, can view the video stream at different resolutions or qualities due to the stream's scalability and embedded properties. For lower bandwidth connections, selective enhancements through bitplane shifting and coefficient weighting can improve visual quality by prioritizing enhancements of certain regions of the video first [2][3]. This approach also can

allow for a picture-in-picture format without doubling the bandwidth. In VoD or PVR applications, the base layer can provide a quick continuous search capability within the bandwidth constraints of the connection. Since elementary streams do not have to come from the same source (this is determined by the BIFS information), enhancement layers can be added to increase the visual quality of an advertisement without a corresponding increase in bandwidth.

Current Streaming Media Technology

Existing media streaming codecs largely used on today's PCs do not take advantage of all the interactive services enabled by a sequence mix of synthetic and natural objects. Current implementation by several equally popular proprietary codecs limit network streaming application to video- and audio-focused material delivered to the PC. As coverage extends to other CPE devices, and implementation develops more personal interactive applications, the complexity in codec and system development will either lead to a limit in growth due to the support of too many proprietary formats or will drive the acceptance of a single format [4]. This single format may arise from a proprietary implementation or from the MPEG-4 standardization process. In either case, the support of the full scope of streaming media on a large scale will require the adoption of MPEG-4 features, if not the exact approaches to implement them.

AD INSERTION

Cable's Response to Existing Local Analog Ad Insertion

Existing ad insertion systems used in today's local cable headends are, at best, hybrid in nature. Locally generated commercials and short programs are stored in an ad server in a digitally compressed format. The digitally compressed commercial is then decoded and converted to analog before insertion into a network program using network cue-tones for timing and duration information of the avail [9], as well as analog splicing techniques [8] .

The hybrid method is fine if content is delivered from a local headend in analog format. But in an all digital delivery environment, a standardized digital-insertion method [9], where a compressed commercial is inserted into a compressed network program in the headend, is the most desirable. With this objective, the SCTE digital video subcommittee (DVS) developed a standard [11] entitled "Digital Program Insertion Cueing Message for Cable (DVS 253)." This standard defines just the cue-message and does not impose any constraints on insertion/splicing equipment. The cue message carries timing information using the coordinated universal time (GPS UTC) for scheduling and MPEG PTS time for frame-accurate insertion, which the splicing device may use to perform the splice. Cue messaging, if required, may be passed on to authorized downstream equipment, such as a pass-through via a remultiplexer to a set-top box. The timing correction needed after remultiplexing also may be transmitted to maintain timing accuracy of the cue signal.

In addition to these two message components, each splice command enables splicing of complete programs or individual components of a program (such as video or audio or data) through the use of `component_tags` enabled by the `stream_id_descriptor`. Today's systems use only program-level splicing, where all the components of a program are replaced at the splice point. In the future, programs and advertisements will be "enhanced" using data broadcast and interactive elements. Typical enhanced commercials could include delivery of discount coupons, prize drawings, or free software. These systems will use component-level splicing, where only selected

components of a program are replaced at the splice point. This splicing enables pre-loading of data streams that are part of the same program by inserting the data component ahead of A/V content. This may be done to load and run the data enhancements in the receiver's application engine.

In addition, the SCTE sub-committee also has developed draft standard DVS 380 [12], which standardizes the API's between the ad-server and splicing equipment. DVS 253, along with standardized APIs defined in DVS 380, will allow splicers and commercial ad servers to interoperate with each other.

Figure 3 shows the functional block diagram of a digital headend where local commercials are inserted into a program utilizing the cue message multiplexed in the transmitted program. The satellite-integrated receiver demodulates the RF signal to its baseband MPEG-2 transport stream and decrypts it. The cue message, multiplexed in the transport stream is detected and is processed by the splicer and ad server. At the splice_insert time, the splicer switches from the network program to the local ad. At the end of the avail, the splicer switches back to the network program. A headend may have multiple splicers and ad servers interconnected among them so that the same commercial need not be stored in more than one server. The cue message standard, DVS 253, does not specify how to splice between two bitstreams. The techniques and resultant constraints for splicing compressed streams have been left to splicing equipment manufacturers. This adaptability also allows for incorporation of advancements in knowledge, such as MPEG-4-enabled technologies.

Digital Splicing into Stat-Mux Channels

In a digital environment, statistical multiplexing of compressed video streams helps to utilize digital channels better than constant bitrate encoded streams, assuming that the peak demand for bits from all video encoders does not coincide. The important constraint here is that all the videos need to be encoded while multiplexing. A previously compressed stream typically had to be decompressed and re-encoded before it could be inserted in a statistical multiplex. This is a detriment to local ad insertion due to the equipment cost of decompression and recompression of the stream as well as visual quality degradation. As an

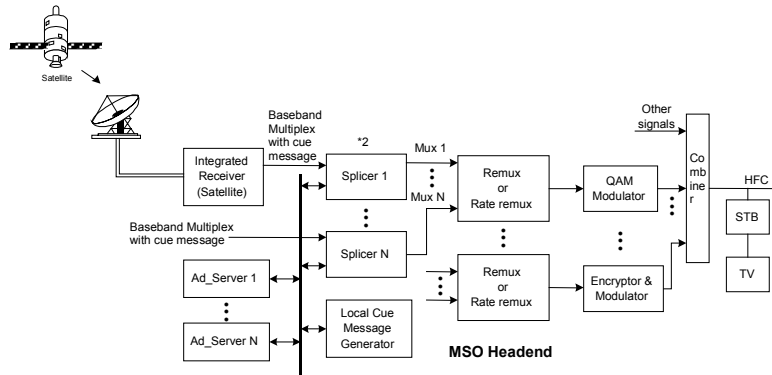


Figure 3. Typical Commercial Insertion System per DVS 253

alternative, a remultiplexer can be used to manipulate or “groom” the multiplexed stream in the compressed domain. By definition, a *remultiplexer* receives one or more multiplexed streams as input and creates a new output multiplexed stream from local operator-selected programs, such as local ads. Nominally, a remultiplexer does not alter bitrates while constructing a new multiplex out of the input streams. The technology that deals with multiplexing compressed video streams and trims the resulting multiplex to match an assigned constant total transmission channel bitrate, is known as *rate-remultiplexing*. Rate remultiplexing meets the latter constraint by transcoding individual video streams within the output multiplex. Transcoding is the technique by which a compressed video stream is translated to a lower bitrate strictly within the compressed domain. It can reduce the bitrate of MPEG-2-compressed video without fully decoding and re-encoding a bitstream. Thus, without cascaded compression, degradation in picture quality is not noticeable with occasional or moderate reductions in the average bitrate of individual video streams.

In addition through rate-remultiplexing techniques, storage requirements for commercials can be reduced by storing only one high-quality version in the server and using rate-remultiplexing to adjust its bitrate. Rate-remultiplexing technology provides the capability to insert compressed digital commercials into digital channels at the headend and removes the need to match the bitrate of the locally compressed commercial with that of the remotely transmitted program, or creating and storing different bitrate versions of the commercials in the ad server.

Targeted Insertion Systems using MPEG-4

While insertion of MPEG-2-coded advertisements into a channel with stat-muxed video requires bitrate transcoding of the advertisement stream, MPEG-4 coding schemes not only allow use of coding tools that are more attractive for advertisements (such as natural video merged with synthetic video with lots of scene changes), they allow for lowering of bit rates to values that are far lower than MPEG-2 stat-mux stream rates. Using a combination of synthetic and natural hybrid tools, advertisements could be authored at rates lower than 500 Kbps with quality that matches that of the MPEG-2-coded video. MPEG-4 provides tools, such as simple, advanced simple, and FGS profiles (visual) capability, and animated 2d mesh, basic animated texture, scalable texture and simple face for coding advertisements on television today.

As advertisements are typically not coded in real-time, MPEG-4’s combination of natural video (including sprites) and synthetic video tools provide a very efficient coding scheme for these applications, even without the use of MPEG-4 system tools, such as BIFS. With the use of MPEG-4 coding, multiple advertisement streams can now be inserted into a stat-mux channel instead of a bitrate-transcoded single advertisement stream using MPEG-2. Currently, standards have been completed in both WG11 for carriage of MPEG-4 streams in networks that carry MPEG-2 transport [1], and Advanced Television Systems Committee (ATSC) that provide for implementation of “targeted advertisements” at the consumer premises equipment by use of appropriate signaling in broadcast streams [11].

Amendment 7 to MPEG-2 systems [1] specifications specifies the insertion of MPEG-4

audio and visual elementary streams into an MPEG-2 transport multiplex with synchronization using the MPEG-2 STC (System Time Clock). In addition, the amendment specifies the use of a MPEG-4 system layer, such as Synchronization Layer (that duplicates some of the PES layer functions). Applications such as AICI (Advanced Interactive Content Initiative) use this amendment for insertion and synchronization of MPEG-4 content with MPEG-2 content. In addition, they also have used BIFS streams to generate advanced compositions on the display screen with user interaction enabled even in a broadcast environment.

The ATSC specification defines signaling based on its system information part of the standard (called Program and System Information Protocol PSIP) for a function called "Directed Channel Change DCC." This is a "virtual" re-tuning of the viewer channel to another part of the transport multiplex at specified timing based on events, as well as user preference settings in the receiver. The switch from the channel being watched to a "directed" channel can occur for criterion such as program identifier, demographic category, postal codes, content subject category, authorization level (premium subscribers) and content advisory values. This information is sent as an MPEG private_section at regular intervals and at the activation time (which could be the cue time of ad-insertion), the receiver can switch to the audio, video and data PID's based on the directed channel change table and switch criterion. Switching back to the original viewed program occurs at the end of the ad-insertion event.

There are two methods of "targeted" ad-insertion that can occur using the above two standards. In the first method, the multiple MPEG-4 based ad-streams are inserted into the stat-mux channel during the ad-insertion period and the DCC directs the receiver to appropriate ad-channel based on user preferences. In the second method, several MPEG-4-based ad-streams can be generated in a multi-program transport stream that can be shared between several stat-mux channels and the DCC for each of the stat-mux channel can direct the user to one of the ad-streams in the large multiplex. The second method allows for larger targeting of streams and for more user preference categories.

Conclusion

This paper discusses some of the advanced features of MPEG-4 standards, which can be implemented to realize advanced A/V content delivery system and targeted commercial insertion in an all-digital environment. Streaming media applications are one of such delivery systems which utilizes network capability effectively and delivers viewer-preferred content on one-to-one basis. Besides, streaming media, which delivers content using IP transport provides unlimited choice to consumers in terms of number of contents available anywhere in the world. This paper also discusses the enabling of low bit rate "effective" advertisement authoring by the emerging MPEG-4 natural and synthetic video coding tools as well as the use of two additional standards that enable "targeted" ad-insertion.

Acknowledgment

The authors would like to thank the management of both CableLabs and Motorola Broadband Communication Sector for their support and encouragement in performing this work.

References

- [1] ISO/IEC 13818-1 (Systems), ISO/IEC 13818-2 (Video) and ISO/IEC 13818-1 Amendment 7, December 2000, ISO/IEC 14696-1 (Systems), ISO/IEC 14696-2 (Video).
- [2] Hyder M. Radha, Mihaela van der Schaar, and Yingwei Chen, "The MPEG-4 Fine-Grained Scalable Video Coding Method for Multimedia Streaming over IP" IEEE trans. On Multimedia, Vol. 3, No.1, March 2001.
- [3] ISO/IEC 14696-2 MPEG-4 Video Amendment 4: Streaming Video profile, MPEG output document N3518.
- [4] Mpeg-4 Overview & Profiles, MPEG home page (www.cselt.it/mpeg/)
- [5] Mukta I. Kar, and Sam Narasimhan, "Targeted advertisement insertion using MPEG-4 coding and SCTE standards for cue-messaging (DVS 253) and API (DVS 380)"

- [6] Mukta L, Kar, Bill Kostka, Majid Chelehmal, and Munsir Haque, "Streaming Over HFC-MPEG-2 or IP or Both?" NCTA paper, 2000.
- [7] Julien Signes, " Binary Format for Scene (BIFS): Combining MPEG-4 media to build rich multimedia services", France Telecom R&D Document, CA, USA 1(650) 875-1516.
- [8] Digital Program Insertion For Local Advertising, Mukta Kar, Majid Chelehmal, and Richard S. Prodan, NCTA paper, 1998.
- [9] Local Commercial Insertion in the Digital Headend, by Mukta Kar, Sam Narasimhan and Richard S. Prodan, NCTA Technical Papers, 2000.
- [10] Rate-remultiplexing: An Optimum Bandwidth Utilization Technology, Richard S. Prodan, Mukta Kar and Majid Chelehmal, NCTA Technical Paper, 1999.
- [11] Amendment 1A to ATSC standard A/65A, "Program and System Information Protocol for Terrestrial and Broadcast and Cable", May 31, 2000.
- [12] Digital Program Insertion Cueing Message for Cable, SCTE Standard DVS 253, December, 1999. Digital Program Insertion Splicing API, SCTE Standard DVS 380 (committee draft).

ISBN 0-940272-01-6; 0-940272-08-3; 0-940272-10-5; 0-940272-11-3; 0-940272-12-1; 0-940272-14-8; 0-940272-15-6; 0-940272-16-4; 0-940272-18-0; 0-940272-19-9; 0-940272-20-2; 0-940272-21-0; 0-940272-22-9; 0-940272-23-7; 0-940272-24-5; 0-940272-25-3; 0-940272-26-1; 0-940272-27-X; 0-940272-28-8; 0-940272-29-6; 0-940272-32-6; 0-940272-33-4; 0-940272-34-2; 0-940272-35-0; 0-940272-36-9; 0-940272-37-7; 0-940272-38-5; 0-940272-39-3; 0-940272-40-7; 0-940272-41-5; 0-940272-42-3; 0-940272-43-1; 0-940272-44-X; 0-940272-45-8; 0-940272-46-6; 0-940272-47-4; 0-940272-48-2; 0-940272-49-0; 0-940272-50-4; 0-940272-51-2; 0-940272-52-0; 0-940272-53-9; 0-940272-54-7

© 2015 National Cable and Telecommunications Association. All Rights Reserved.