

THE COMPLETE
TECHNICAL PAPER PROCEEDINGS
FROM:



AN EVALUATION OF ALTERNATIVE TECHNOLOGIES FOR INCREASING NETWORK INFORMATION CAPACITY

Ron Shani, Xtend Networks

David Large, Consultant

Abstract

HDTV, VOD, ITV and other applications are placing ever-greater pressure on operators to transport more information – that which is widely distributed as well as communications with individual customers. Choosing how to create adequate capacity is difficult; driven by financial and regulatory constraints, capital costs and ongoing operating considerations.

This paper will evaluate some of the technical options against those factors. Evaluated technologies will include bandwidth expansion to 1 GHz, more efficient modulation, more efficient video encoding, elimination of analog video carriage, splitting of existing nodes, switched digital video and a proposed use of frequencies above 1 GHz that offers the greatest bi-directional bandwidth expansion and the greatest benefit/cost ratio.

INTRODUCTION

Bandwidth Pressures

The history of cable television is one of ever-increasing need for information capacity, initially driven by the expansion of over-air broadcasting, then premium and ad-supported satellite networks, followed by pay-per-view and high-speed data services. Today, operators are launching bandwidth-intensive high-definition television (HDTV) channels, various flavors of video on demand (VOD), higher Internet access data rates, and telephone services. Each of these increases the need for an increase in system information capacity (for purposes of this paper, unless otherwise specified, “bandwidth” will be used

interchangeably with “information capacity” and “RF bandwidth” will be used when the historical meaning is intended). The increasing bandwidth demands fall into three broad categories:

1) Common downstream (“broadcast”) bandwidth; that is, bandwidth occupied by signals that are transmitted throughout the network (irrespective of whether or not individual customers are enabled to receive them). An example of common signals would be a high-definition stream from HBO that would be continuously transmitted system-wide, but for which only certain subscribers would be authorized.

2) Interactive downstream (“unicast”) bandwidth; that is, bandwidth occupied by signals that are transmitted to individual customers. VOD, Internet communications and telephone are all examples of such signals.

3) Upstream bandwidth; that is, bandwidth occupied by signals that are transmitted from individual customers towards the headend. With the exception of a small amount of bandwidth occupied by network element management systems (NEMS), all upstream signals fall in the same category as interactive downstream bandwidth.

The Case for Dramatic Bandwidth Increase

Historically, manufacturers have offered cable operators increases in upper downstream RF bandwidth limits in steps of 50 MHz or so from an upper frequency limit of 220 MHz to 860 MHz, with the upstream bandwidth remaining fixed, except for one step from 30 to 42 MHz. By contrast, in the

data world, speeds have increased exponentially over several orders of magnitude. As more content carried over cable systems is digital in nature, more communications are directed to and from individuals, and competitors greatly increase both video and non-video capacity, the question is whether operators will need to significantly increase bandwidth, especially upstream bandwidth, to take advantage of opportunities and meet competition.

A few points to consider:

1) On the competitive data front, SBC and Verizon, among other telcos, have launched a major fiber-to-the-curb/home push. Typical of the technology to be deployed is Wave7's equipment which provides 500 Mb/s symmetrical data, shared among 16 passings, in addition to 860 MHz of RF downstream bandwidth.ⁱ Verizon is offering data rates to 30 Mb/s downstream/5 Mb/s upstream in its fibered markets, with the capability to offer rates of hundreds of megabits per second.ⁱⁱ Some overbuild competitors in the US have already offered 100 Mb/s service options to customers and speeds of between 10 and 100 Mb/s are commonly available in Asia. Finally, the capability of copper plant continues to improve and now supports high-definition digital video.

2) Cable operators are already being pushed to significantly increase rates – Comcast announced a standard rate of 4 Mb/s and an available 6-Mb/s downstream/768-kb/s upstream rate; Cox increased its standard rate to 4 Mb/sⁱⁱⁱ and RCN has upgraded its rates to 10 Mb/s.

3) On the telephone side, the number of VoIP residential and small business lines is predicted to hit almost 11 million by 2008, with a significant amount of that traffic carried over cable systems.

4) Direct broadcast satellite operators will be taking advantage of new spectrum, closer satellite spacing, higher power and spot beam technology to realize greatly increased throughput – as much as 18,000 MB, or enough to carry 2800 high-definition programs at 6.5 Mb/s/program using advanced codecs.^{iv}

5) In general, television is moving from pre-scheduled broadcast of standard-resolution programs to on-demand presentation of high-definition, with a 4X increase in bits per stream and the need to send programming to (and receive communications from) individual subscribers. Competitively, one satellite operator expects to offer its customers 150 national and 500 local HDTV channels by 2007

6) Finally, upstream data communications rates from subscribers are increasing rapidly. VoIP is a symmetrical service; file sharing can be symmetrical or even asymmetrical in the upstream direction; and near-future services such as video telephony will require multiples of the bandwidth required for voice. Comcast recently announced plans to offer video instant messaging. RCN now offers a video surveillance service that allows customers to stream video from up to four cameras through their broadband connection.^v

In summary, there is significant evidence that cable operators will need major increases in bi-directional information capacity in the near future, and that the upstream in particular, with a current capacity of only about 100 Mb/s/node, is a major bottleneck that will need to be addressed.

Operators can realize this increased information capacity through an increase in RF bandwidth, through more efficient use of existing bandwidth, or through more efficient sharing of existing bandwidth. Additionally, increased interactive bandwidth can be realized by sharing of the bandwidth devoted

to interactive services among fewer customers. The various upgrade technologies that will be considered differ in their effects on broadcast versus interactive and downstream versus upstream information capacity, as will be seen.

Candidate Technologies

There are many approaches to generating more information capacity in a cable system. This paper will evaluate the following possibilities:

- 1) An increase in downstream upper RF bandwidth limit from 550, 750 or 870 MHz to 1 GHz.
- 2) An increase in digital modulation density from 256 QAM to 1024 QAM.
- 3) Utilization of more effective digital video compression technologies, such as MPEG-4.
- 4) Subdivision of existing optical nodes.
- 5) Elimination of analog video carriage, with the formerly-analog signals transmitted only in digital form.
- 6) Use of switched digital video to avoid sending low-usage channels to subscriber groups except when requested.
- 7) Use of RF bandwidth above 1 GHz to expand both downstream and upstream capacity.

This is obviously not a comprehensive list, and the choices are not mutually exclusive. For example, an operator may choose to simultaneously increase modulation density and also use advanced digital compression algorithms. For keep the matrix manageable, however, we evaluated each option separately.

When it comes to discussing quantitative results, we used what we felt were reasonable assumptions for an average cable system. For every possible upgrade scenario, however, the results will vary depending on the assumed condition of the unmodified plant. For example, a marginal system may not be able to take advantage of 1024 QAM without

fixing basic problems, while other systems may require little incidental preparation.

Methodology

For ease of comparison, each technology was evaluated as a candidate for upgrading a hypothetical 100,000 home cable system which currently has 500-home nodes and an average density of 100 homes per plant mile. It is assumed to be 80% aerial plant. The connected household penetration is assumed to be 70%, with 35% of connected homes equipped for digital video reception. The system is assumed to currently carry 80 channels of analog video, 136 total standard-resolution broadcast digital video streams, 12 high-definition broadcast digital video streams, VOD, high-speed data, and VoIP. Unless otherwise stated, the system is assumed to have been upgraded to 750 MHz within the previous ten years. Other assumptions regarding the system will be discussed when relevant to each individual candidate technology.

Technologies were evaluated with respect to their effect on both downstream and upstream capacities and with respect to both commonly delivered (broadcast) and interactive services. In each case, the technologies were also evaluated qualitatively with respect to future enhancement options. Finally, conformance of each alternative to current regulatory requirements is noted.

INCREASE TO 1 GHz BANDWIDTH

An increase in the upper downstream frequency limit to 1 GHz follows the traditional pattern of cable RF bandwidth expansion. While it offers additional downstream capacity, it does not address the upstream bottleneck and does not offer a straightforward path to future increases, as discussed below.

Distribution Network Issues

The cost of coaxial equipment upgrade will depend on the starting bandwidth and on the condition of the original plant. Variables include: the percentage of passive devices which are already rated at 1 GHz, whether upgrade modules are available for actives, whether the gain of the new actives will be sufficient to avoid re-spacing, the condition and type of original coaxial cable and connectors, and whether the increase in drop cable loss is such as to require replacement.

The tradeoffs in a bandwidth increase are well known. If the amplifier spacing does not change, each amplifier must have higher gain and either the input levels will be lower (degrading C/N), the output levels must be higher (degrading distortions), or the amplifier must have higher power output hybrids (increasing power consumption and heat). Furthermore, the number of signals carried will presumably increase, further increasing intermodulation products. Alternately, amplifier spacing can be decreased, but then the number of cascaded amplifiers increases, degrading both noise and distortion. Thus, this technology is self-limiting and does not offer a solution to future expansion.

Our estimates are based on figures developed by a major MSO for their current mix of cable systems of various bandwidths, conditions and original parentage. Added to these costs are estimates of the replacement optical equipment required at headend or hub and node to feed the upgraded plant.

Consumer Premises Equipment (CPE) and Regulatory Issues

The entire cost is not in the distribution system upgrade – the bandwidth must be used for something. No existing CPE tunes above 870 MHz, nor is it required to do so to meet current DOCSIS (data) or SCTE 40 (video)

standards. Furthermore, since the FCC has adopted SCTE 40 into its rules, operators are forbidden from offering one-way digital video services above 864 MHz.

We therefore assumed that, while the upgrade would create additional capacity between the existing upper limit and 864 MHz, the space above that would be limited to services that need be received only on CPE provided by cable operators. Of the available choices, the most logical seemed to be simulcasting of the existing analog programming (the first step to an eventual all-digital plant and recovery of the spectrum now used for analog transmission) to digital-only converters. We estimated the cost of simulcasting from a report on Charter's Long Beach, CA conversion^{vi} and estimated the cost of the digital-only converters at \$85^{vii}, the recovery value of the old converters at \$25, and the labor cost to make the change at \$10. Thus, the estimated cost includes the CPE changes necessary before the expanded bandwidth can be used, but not the cost of adding any new services.

Using these assumptions, the total cost and gained downstream bandwidth (in equivalent 6-MHz channels) is as follows:

Original Bandwidth	550	750	860
Cost/HP	\$274	\$116	\$81
Added DS Chans	75	42	23

This upgrade, of course, does nothing to address the upstream issue.

UPGRADE TO 1024 QAM

The highest existing digital modulation is 256 QAM, which transmits 8 bits of information per symbol, for an effective transmission rate of about 38 Mb/s in a 6-MHz RF channel. One proposal for increasing information capacity is to use the next logical increment of modulation density, 1024 QAM, to increase the bandwidth

efficiency of networks by transmitting 10 bits per symbol, a theoretical increase of 25%.

The practical network issue with this upgrade is existing network noise and distortion performance. SCTE 40 mandates end-of-line C/(noise + interference) of 33 dB for 256 QAM.^{viii} To maintain the same headroom, a 1024 QAM signals would need to be received with a C/(noise+interference) of 39 dB.

In most cable systems, data signals are carried at the same average power level as analog video signals (typically referred to as 6 dB lower only because analog video signals are referenced to sync peak level and digital signals to average power level). Thus, raising the power level of 1024 QAM signals is probably not a practical option.

Typical cable systems are designed for an end-of-line ideal analog video C/N (thermal noise only) of 48 dB. With normal variations, aging and maintenance tolerance, 46 dB is about all that can practically be assured – just enough to pass the FCC’s 43 dB requirement after passing through a typical converter (with 0 dBmV input and a 13 dB noise figure).

Taking into account the difference in noise susceptibility bandwidth between video (4 MHz) and data (5.3 MHz) and the 6 dB difference in how their levels are referenced, the expected carrier-to-thermal-noise of a received data signal may be as low as 38.8 dB, to which must be added the effects of composite beat products among analog signals, composite intermodulation products among digital signals and crosstalk in multi-wavelength optical links.^{ix} Otherwise stated, a system that just meets FCC specifications for analog video will not be adequate to carry 1024 QAM signals.

Distribution Network Costs

To account for solving the inevitable system problems and increasing performance

slightly, we estimated a cost of \$850/mile in distribution system “fixes”.

Headend Costs

Existing headend modulators must be replaced to prepare the system to utilize the expanded throughput. To minimize the cost, we assumed that only digital video modulators are replaced (leaving data and VoIP unchanged). Additionally, we added the cost of re-multiplexers for those signals currently received from satellite and passed through the headend unchanged. This prepares the system for adding 2-3 additional video streams per multiplex in the future.

Customer Premises Equipment Costs and Regulatory Issues

No existing CPE is capable of receiving 1024 QAM signals. Furthermore, current FCC rules mandate that one-way digital video services use only 64 QAM or 256 QAM. Although operators may approach the technical issue in various ways, we assumed that existing converters would be replaced with hybrid analog/digital converters enhanced to receive 1024 QAM that would cost \$175, with a value of \$25 assigned to the retrieved converters they replace. We have assumed that all existing converters are replaced, which enables the efficiency improvement to be applied across all digital video channels, but means that converter replacement dominates the other costs.

Using these assumptions, the total cost of an upgrade to 1024 QAM is \$60 per home passed for an effective downstream bandwidth increase of 6.75 6-MHz RF channels. As with the 1 GHz upgrade, converting to 1024 QAM does not address the upstream bandwidth constraint, nor provide a path for future upgrades. Unlike, a 1 GHz upgrade, our 1024 QAM scenario is in conflict with current FCC regulations, however applying it to only

interactive services would greatly reduce the throughput gain.

ADVANCED VIDEO COMPRESSION

All cable digital video services today are compressed using MPEG-2. While this was a breakthrough technology when introduced, more efficient algorithms have since been introduced, of which the dominant contenders are MPEG-4 AVC and Windows Media (SMPTE VC-1). Either offers roughly a 2:1 increase in streams/channel compared with MPEG-2. Since a large use of downstream bandwidth in a typical cable system is for digital video, adopting a more efficient compression algorithm will increase overall effective throughput.

Because no new modulation is involved, the upgrade imposes no increased demand on the distribution network.

Headend Costs

The cost of adopting advanced compression is dependent on how widely it is adopted. We assumed that most digital video would arrive at the upconverted system in the new format, either from the original program source or from the MSO's regional center.

Customer Equipment Costs and Regulatory Issues

The CPE situation for advanced encoding is essentially the same as for use of 1024 QAM – existing boxes do not receive and cannot be upgraded to receive the new-format signals, and thus require replacement.

The regulatory issues are also similar, as the FCC limits one-way digital services to MPEG-2 encoding. As with 1024 QAM, we calculated the efficiency gain across all digital video channels and did not address the regulatory issues.

In summary, an upgrade to advanced video encoding is less expensive than an upgrade to 1024 QAM because no plant changes are required, and results in a larger effective capacity increase. Specifically, the estimated cost, using our assumptions, is \$53 per home passed and results in an effective bandwidth increase of 12.2 downstream RF channels. It does not address the upstream bottleneck.

NODE SUBDIVISION

Subdividing optical node serving areas does not increase the instantaneous system information capacity to any network segment, but does share that capacity among fewer customers. Thus, to the extent that the sub-areas are fed separately, an effective capacity increase is realized for those services which are delivered to individual customers. To be precise, the capacity is increased in proportion to the bandwidth allocated to those services and multiplied by the number of downstream or upstream segments created. Furthermore, since nothing is changed except for effective node size, no regulatory or CPE technical issues are created.

Plant Costs

We assumed that the previous upgrade was not a total rebuild – that is, it utilized as much of the then-existing plant as possible -- and that the cost was further minimized by “dropping” non-scaleable nodes into the coaxial distribution system to create the required 500-home serving areas. Thus, the cost of node subdivision included the cost of replacing the node itself with a segmented model (2:1 downstream and 4:1 upstream) and re-routing the coaxial distribution plant to create four roughly-equal-sized segments (requiring, on average, 1,000 ft of new cable plus splicing). It does not include any service-specific hardware.

Headend Costs

In order to activate the expanded bandwidth, we included the cost of one additional downstream transmitter and three additional upstream receivers to communicate with the new sub-nodes and thus activate the additional capacity.

In summary, we estimated that the division of 500-home nodes into two downstream segments and four upstream segments, would cost approximately \$30 per home passed. We assumed that eight downstream channels were used for individual subscriber, interactive services, resulting in a net effective bandwidth gain of four channels in each of the two downstream sub-nodes, equivalent to a doubling of the downstream interactive service throughput capability. We assumed that 30 MHz of the upstream bandwidth was usable for interactive services and therefore the 4:1 split creates an effective bandwidth gain 22.5 MHz in each of the four upstream sub-nodes, equivalent to a quadrupling of upstream interactive service throughput capability.

CONVERSION TO ALL-DIGITAL

In the future, all television, whether over-air broadcast, satellite or locally originated, will be in digital form. One option for operators is to accelerate that process by converting all current analog video signals to digital form and providing digital converters at every connected television receiver.

The advantages include at least a 10:1 increased usage of former-analog bandwidth, lower cost receivers, uniform transport protocols across all services and breaking DBS operators claim to be the only “all digital” network. Disadvantages include the cost of providing converters to current analog subscribers and defeating the features of some basic subscriber’s video equipment. Additionally, current FCC regulations require

carriage of at least Basic channels in analog format absent individually-granted exceptions, though that requirement will cease when broadcasting transitions to digital^x.

We evaluated two versions of an all-digital conversion – a downstream-only version and a further option in which a portion of the formerly-downstream bandwidth is allocated to upstream usage.

Plant Costs

Since standard 256 QAM signals are assumed, analog channels are converted to digital at approximately the same total RF power per channel, and thus no additional loading is placed on the distribution system. As discussed above, end-of-line digital signal C/N should be slightly below 39 dB at worst, and thus have a significant margin above the SCTE 40 and FCC minimum of 33 dB, even when distortion parameters are included. Thus, no plant changes are required to make the analog to video conversion in the downstream-only option.

Expanding the upstream bandwidth, however, requires changing every duplex filter in the system, the upstream amplifiers (wider bandwidth and higher gain), upstream optical transmitter modules in nodes, and optical receivers in the headend. We assumed that the new upstream spectrum would extend from 10 to 85 MHz and that the downstream spectrum would start at 105 MHz to preserve use of equipment that operates in or near the FM band. We estimated the total of plant and optical headend cost to make the frequency change to be \$10,180 per 500 HP node, including the cost to realign the plant.

Headend Processing Costs

As with advanced compression techniques, we assumed that most (75%) of signals would arrive at the headend in digital form from broadcasters, cable networks or MSO regional

centers, but that the remainder would require conversion at the headend. We scaled Charter's reported cost to upgrade their California system^{xi} by the required number of locally-converted channels and estimated the total headend cost to be \$250,000.

Consumer Premise Costs

We assumed the same \$75 digital-only converter cost for this option as for the 1024 QAM case. The difference is that, rather than replacing existing digital converters because of incompatibility, additional converters are required for every television outlet in the system what did not previously have one.

The cost of the downstream-only ("low-split") version and the version that includes expanding the upstream spectrum ("mid-split") is summarized in the table below. The mid-split version more than doubles upstream capacity.

Option	DS Only	DS + US
Cost/HP	\$157	\$175
Added DS Chans	72	63
Added US MHz	0	38

SWITCHED DIGITAL VIDEO

With the exception of server-based on-demand programming, cable operators currently transmit all available programming choices simultaneously and continuously throughout their networks. However, given the widely different popularity of different programming among any given group of subscribers, viewing is concentrated among a few channels and many of the hundreds offered are not simultaneously viewed. Thus, even though there are good reasons for offering a wide choice of programming, it is an inefficient use of bandwidth to send signals to sections of the network except when at least one subscriber wishes to access them.

Switched digital video (SDV) gains effective network throughput by offering less popular programs to service groups only on demand, using technology that is transparent to users – that is, the viewer should ideally be unaware when selecting a program that it might not be delivered until the virtual channel is selected. Use of SDV does not increase the information capacity of the network, but rather shares it more efficiently.

When a switched channel is selected, a small resident application in the user's box sends a request to the headend SDV server, which, if the requested stream is not already being viewed in the service group, adds it to an appropriate multiplex. Then it directs the box to the correct channel and program identifier. When the channel is no longer being viewed within the group, the stream is dropped.^{xii}

Trials of SDV are still in an early stage, with widely varying results. One operator estimated a potential savings of 26% of total video channels^{xiii}, while a larger and more recent trial conducted in a Cox system suggests that as many as 41 programs can share an RF channel on a switched basis and that trials with 28 programs per channel resulted in no instances of blocked access^{xiv}.

The regulatory problem with switched video is that operators are required to deliver all non-interactive digital video services in a way that is compatible with one-way digital cable-ready receivers. Those receivers are obviously not capable of sending message to the headend to request streams. Thus, until that hurdle is overcome, only interactive video services can be offered on a switched basis.

Given the state of development of SDV and regulatory constraints that limit which channels can be offered on a switched basis, we assumed that 100 current two-way digital program offerings (premium and pay-per-view) would be delivered over five

statistically-shared RF channels – a savings of 50% over the spectrum formerly required. In other words, we assumed that an operator would choose in this option to comply with current regulations.

Headend Costs

There are no distribution plant costs associated with the addition of SDV, since the distributed signals are identical to non-switched signals. In the headend, a SDV manager is required to manage the addition and deletion of streams and an MPEG switch/mux is required to create the required multiplexes feeding each service group. We assumed that bandwidth, multiplexers and modulators were dedicated to the SDV service, rather than being shared with on-demand services. The cost estimate assumes 4 nodes and 5 RF channels (80 streams total) per service group.

Customer Premise Costs

SDV is completely compatible with current-generation digital two-way boxes. The required software module is much smaller than that required to implement VOD. Therefore, there is no cost to implement SDV among existing digital video subscribers.

With the above assumptions, implementation of a limited SDV service to existing digital video customers is estimated to cost only \$5175 per node, but to free up only the equivalent of 5 downstream RF channels. It has no effect, of course, on upstream congestion and, in fact, adds traffic from video set-top boxes for which low latency is very important. Much greater gains are possible, of course, but only if the regulatory issues are resolved.

EXTENDED BANDWIDTH (> 1 GHz)

A final choice is to activate the spectrum above 1 GHz for bi-direction bandwidth

expansion. While various proposals for use of this spectrum have been proposed for many years, none have been widely deployed as an upgrade strategy. The system described below, however, has been successfully used to implement selective overlays to service commercial customers^{xv}, so the viability of the technology was not at question, but rather its applicability to a system-wide upgrade to serve the entire customer base.

The version we evaluated is based on the creation of two sub-octave transmission bands -- 1250-1950 MHz downstream and 2250-2750 MHz upstream -- with nodes and amplification equipment paralleling legacy equipment in the field, as shown in Figure 1. Splitting amplification between legacy and extended amplifiers greatly simplifies amplifier design. Passives are replaced by equivalent units passing the entire 5-2750 MHz band. Tests have shown that current-generation hard cables used by the industry will support these frequencies, while the first non-TEM mode does not occur until about 4.6 GHz for the largest cable sizes in use today.^{xvi}

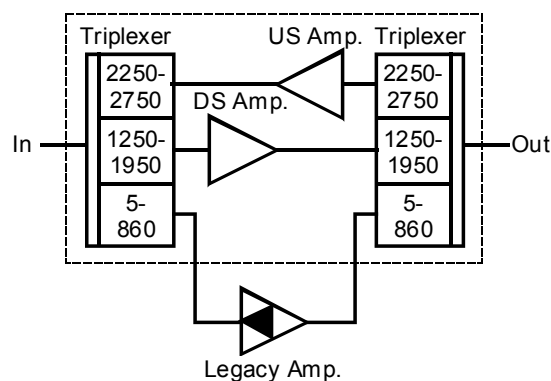


Figure 1: Extended Frequency Amplifier

In order to assure compatibility with existing headend and CPE and to ensure compliance with FCC regulations, the extended downstream frequency band is converted to/from 100-800 MHz and, for residential applications, each upstream sub-block of 12-42 MHz could be converted to one of ten “slots” in the 2250-2750 upstream distribution band and back in the headend.

This scheme would allow both headend RF equipment and consumer premises equipment to operate at normal levels and frequencies. Since the entire 500 MHz upstream spectrum is transported to the headend without any channelization, it is also possible to allocate a wider portion of the spectrum to applications which require greater data rates than can be transported through a standard cable television upstream path. Only a terminal equipment change would be required to re-allocate the spectrum for such applications. As with any block segment conversion scheme (such as those used for years for upstream node segmentation) low phase noise, frequency accurate converters are required. The premise equipment diagram is shown in Figure 2.

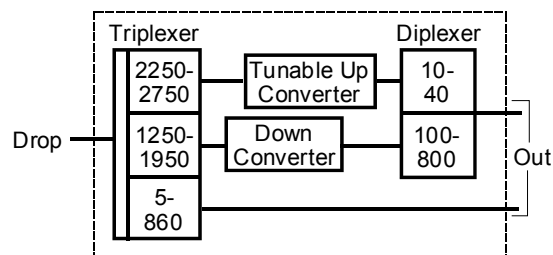


Figure 2: Premises Terminal

In addition to a 2:1 downstream and up to 17:1 upstream bandwidth increase (11:1 if all channelized in 30-MHz sub-bands), a significant advantage of this scheme is that the expanded upstream communications capacity is not “locked” to a sub-node group (as in node subdivision), but rather can be assigned on a subscriber-by-subscriber basis anywhere in the node serving area, simply by assigning which slot upstream frequencies are converted to. A second significant advantage is that the ingress into any upstream slot is limited to that occurring in the residences (or businesses) whose upstream communications are converted to that slot, thus improving the usability of the added spectrum. Ingress in the drop has no effect on the extended spectrum.

Plant Costs

The plant upgrade effort is comparable to a typical rebuild effort. Because the required techniques to parallel active equipment are different, we based our estimated labor costs on actual field trials and on evaluations by experienced contractors in cable construction. The material costs include powering upgrades required for the additional actives. As with node subdivision, we assumed that the required two fibers (or at least one wavelength on each of two fibers) between headend or hub and node were available.

Headend Costs

Headend costs include the equipment required to activate the additional bandwidth, including frequency conversion and optical transmitters and receivers.

Premise Costs

As with other options, the cost of activating frequencies above 1 GHz is dependent on how those frequencies are utilized. For purposes of this study, we assumed that analog video, cable modem and VoIP signals remained on the legacy bandwidth, while digital video and a digital simulcast of analog video signals were placed on the expanded downstream band, with STB upstream signals on one of the expanded upstream slots. This scenario is compatible with all existing digital and analog equipment, including subscriber-owned modems and television receivers, while allowing the cable operator to purchase digital-only STBs going forward and freeing bandwidth for advanced video and data services. The model included the cost of installing residential block converters for every customer who subscribes to digital video services. We estimated the labor cost of this to be comparable to installing a standard drop amplifier.

We evaluated the cost of such an upgrade under two scenarios -- activation of two or eight of the ten possible upstream slots -- in order to determine how sensitive the technology is to the degree of upstream bandwidth expansion. The results are summarized in the following table.

Option	2 Blocks	8 Blocks
Cost/HP	\$121	\$129
Added DS Chans	117	117
Added US MHz	60	240

In summary, the use of frequencies above 1 GHz for expansion of both down and upstream bandwidth offers the greatest information capacity of any of the evaluated options, with the possibility of further expansion of the upstream at very low cost.

SUMMARY

Quantitative comparisons will depend on assumptions and intended use the expanded capacity. This summary is based on the assumptions stated previously.

Regulatory Issues

The use of frequencies above 862 MHz, 1024 QAM or advanced encoding for digital video all violate provisions of Paragraph 76.640 of the FCC's rules, if applied to one-way digital video services. Until and unless those provisions are modified, the gain from use of these techniques will be constrained because of that. Our summary results take those restrictions into account.

Secondly, the FCC requires that basic television service be carried in an analog form. Thus, the conversion to all-digital video is dependent on obtaining a waiver or waiting until all VHF over-air transmission ceases.

Instantaneous vs Virtual Capacity Increases

Some evaluated technologies increase the peak information-carrying capacity of the network, while others realize the effective throughput increase by other means. Both are important: Peak capacity limits the amount of information that can be transmitted to any given subscriber group, while virtual capacity increases are dependent on how services are divided between those which are broadcast and those which are directed to specific customers or customer groups. Today, peak upstream capacity, using 16 QAM, is limited to about 100 Mb/s (ten 3.2 Mb/s channels, each with a capacity of 10 Mb/s). Even if 64 QAM were usable across the entire 9-41 MHz band, the potential increase would only be about 25% to 125 Mb/s.

All the evaluated technologies increase downstream effective capacity. The following table shows which also increase peak downstream information rates and which increase upstream effective and/or peak rates.

Technology	DS	Upstream	
	Peak	Peak	Virtual
1 GHz	Yes	No	No
1024 QAM	Yes	No	No
AVC*	Yes	No	No
Node Split	No	No	Yes
All digital	Yes	No	No
+ US expand	Yes	Yes	Yes
Switched	No	No	No
Extended BW	Yes	Yes	Yes

*Advanced video compression

Only the elimination of analog video combined with expansion of the upstream band or the use of two-way extended bandwidths provides an increase in both upstream and downstream effective and instantaneous information rates.

Comparisons of Capacity and Cost

Figure 3 illustrates the increase in effective downstream channels for each of the technologies, while Figure 4 shows the increase in effective upstream bandwidth. Figure 5 shows the cost effectiveness of each, which we calculated by taking the ratio of per-node capital cost to the total downstream-plus upstream effective bandwidth increase.

Assuming our assumptions are reasonable, it appears that the most efficient 1 GHz upgrade is from a 750 MHz system. While a 550 MHz system will gain more DS bandwidth, it will also require much more cable, passive and drop replacement work. On the other hand, while an 860 to 1 GHz upgrade is the least costly of the three, the lower incremental bandwidth makes it less efficient.

As expected, converting to all-digital video gains a lot of DS bandwidth due to the 10:1 improvement in program streams per channel. Looking at Figure 5, however, it is not one of the most capital-efficient upgrades simply because of the cost of placing one or more digital converters in every Basic subscriber's house. Converting to a mid-split configuration is slightly less cost-efficient but is one of only three options to improve the critical upstream throughput bottleneck.

1024 QAM, advanced video compression and switched video offer only moderate

throughput gains for a couple of reasons. 1024QAM, which carries 10 bits per symbol, only offers a theoretical 25% gain over 256 QAM. Additionally, all three technologies are currently constrained the FCC regulations and which are derived from SCTE40. Of the three, SDV is the most efficient because it is compatible with existing set-tops. Absent regulatory restrictions, SDV holds the promise for major downstream effective throughput gain.

Splitting of existing nodes offers significant gains in effective downstream and upstream bandwidth for moderate cost and without causing any regulatory problems or equipment compatibility issues. It is second only to extended bandwidth in cost efficiency.

Use of extended bandwidths, as described earlier, is comparable in cost to a 1 GHz upgrade and less expensive than an all-digital conversion. It offers the greatest incremental bandwidth improvement -- effective and instantaneous; upstream as well as downstream -- of any of the options. As a result its cost effectiveness is greater than any of the alternatives. Furthermore, the incremental cost to activate 8 upstream blocks is slight compared with activating just two, so that it is very economically scalable to future expansion needs. We suggest that it should be seriously considered for future major throughput upgrades.

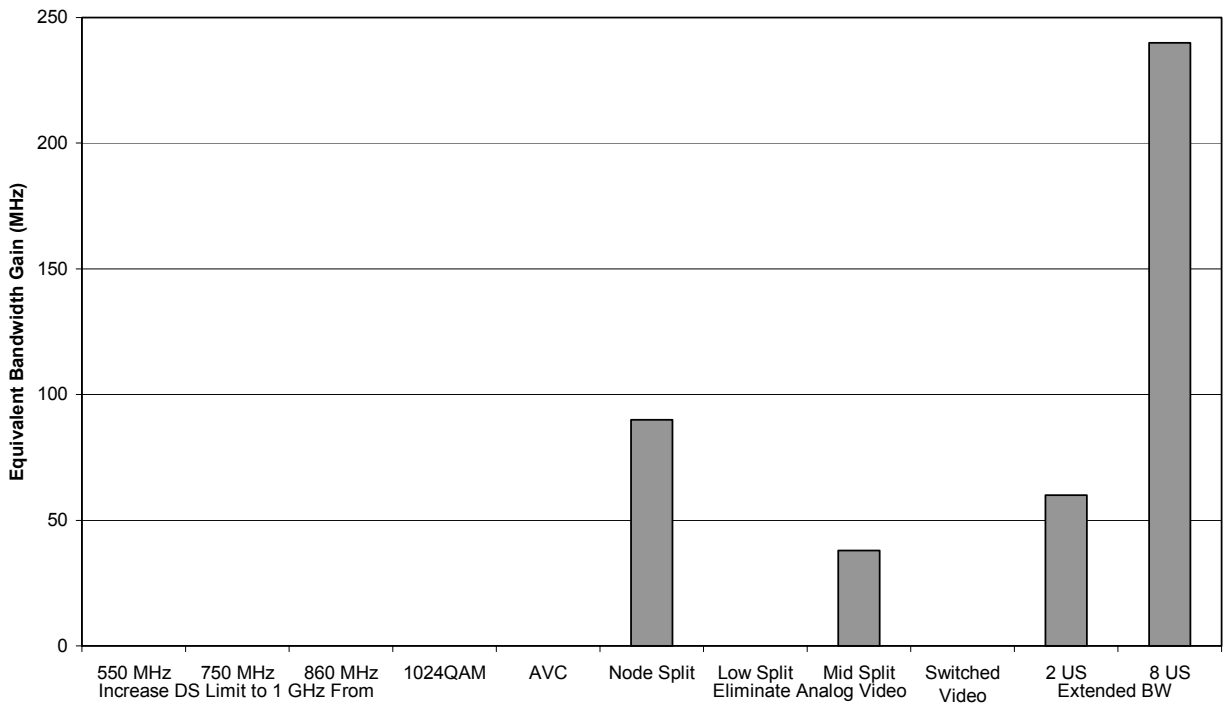


Figure 4: Effective Upstream Bandwidth Gain

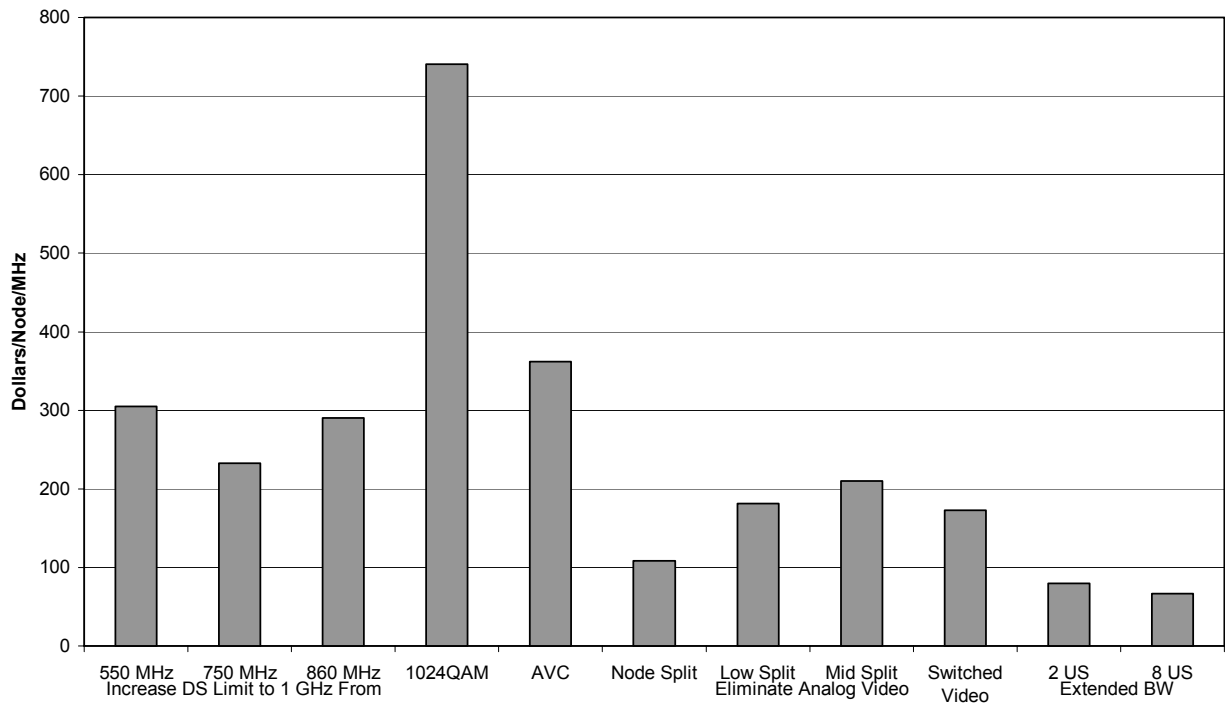


Figure 5: Cost-Benefit Ratio

REFERENCES

1. Report on SCTE's 2005 Emerging Technology conference by Laura Hamilton and Johathan Tombes in Pipeline.
2. "Verizon identifies more 'FiOS' markets," CED Magazine online, 1/19/05.
3. "Linksys adds ADSL2+ support," (describing 25 Mb/s technology) Jeff Baumbartner, CED Broadband Direct, 2/3/05.
4. "Comcast raises broadband speed," Jim Hu, c/net News.com.
5. "RCN pushes cable-modem speed to 10 mbps," CTAM SmartBrief, 1/20/05.
6. "What's ahead for Vonage in 2005?" Net.Worker column in Network World, 1/10/05.
7. "Satellite Technology and Platforms," Steven Osman, presented at SCTE's Emerging Technologies 2005 conference.
8. "Cable could rule if it plays it cards right," David Lieberman, USA Today.com, 1/24/05
9. "Comcast puts video IM on deck," Jeff Baumbartner, CED Broadband Direct, 1/20/05.
10. "RCN offers Net-connected surveillance," Boston Globe, January 27, 2005.
11. "Convergence, California Style," IP Solutions Summit, Supplement to MultiChannel News, CED and Broadcasting and Cable magazines, September 2004.
12. "Coneco offers all-digital option," CED Broadband Direct, 9/20/04. Also, "Bargain boxes, is the cable industry within reach of the \$50 all-digital set-top," CED Magazine, September 2004.
13. SCTE 40 2004 Digital Cable Network Interface Standard, Table B.
14. See Chapter 13 of Modern Cable Television Technology, Second Edition, by Ciciora, Farmer, Large and Adams for a fuller explanation of fiber crosstalk effects.
15. "Congressional leaders want 'hard date' end for analog broadcasts," (suggesting 2006 deadline) CTAM SmartBrief 2/3/05
16. See footnote 8.
17. "Getting from here to anything, anytime," David Large, CED magazine, September, 2004. This article explores and advantages and limitations of SDV.
18. "IP-based switched broadcast," Joachim Vanhaecke, Proceedings Manual, 2003 SCTE Cable-Tec Expo.
19. "The statistics of switched broadcast," Nishith Sinha and Ran Oz, Proceedings of SCTE 2005 Conference on Emerging Technologies.
20. "Arris, Xtend test out-of-band DOCSIS," Jeff Baumbartner, CED Broadband 1/31/05
21. Modern Cable Television Technology, 2nd edition, equation 10.11 (which is in error by a factor of 2).

BANDWIDTH MANAGEMENT FOR THE UNIVERSAL EDGE

By Bruce Thompson, Xiaomei Liu
Cisco Systems, Inc.

Abstract

To offer more services to end users in the cable network, bandwidth needs at the edge of HFC network are growing rapidly. The industry is moving from current architecture where each service has its own edge resources to a multi-service universal edge architecture to use the RF bandwidth more efficiently.

This paper will go into detail on the benefits of dynamically sharing resources across services. It will also describe a control/data plane architecture that can be used to share resources between these services. Lastly, this paper will provide examples of standardized protocol suites that can be used to implement the interface between the components of a distributed architecture that supports resource sharing across services.

OVERVIEW

The HFC plant is the point of convergence for all of the services that MSOs provide. In addition to Docsis based Internet Access and Broadcast Video, new services such as video on demand (VoD), network PVR and switched broadcast are now being offered. In current deployments, each service has a statically allocated portion of the RF spectrum. Dedicated QAM pools are allocated for broadcast, VOD, switched broadcast and DOCSIS based Internet Access services.

Using VOD and DOCSIS Internet Access services as an example, Figure 1 shows the

deployment scenario where QAMs are dedicated to each service. For VOD services, VOD servers send content over gigabit Ethernet (GE) to video capable downstream (DS) QAMs which reach Setop boxes (STB) in the home. VOD servers are in control of allocating both video pumps and QAMs when VoD session requests are made by STBs.

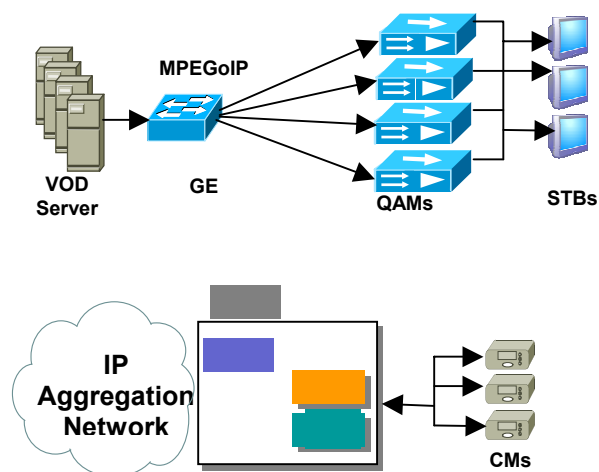


Figure 1. VOD and DOCSIS services in today's cable network

For DOCSIS Internet Access services, a CMTS serves a group of cable modems. The CMTS consists of a Docsis MAC layer processor, downstream (DS QAM) channels, and upstream (US) channels bundled into a single platform. The downstream QAMs embedded in the CMTS use a different portion of the RF spectrum than the QAMs dedicated to the VoD service.

With the increasing popularity of VoD services, high definition television, and DOCSIS Internet Access, the demand for HFC network bandwidth is ever increasing.

MSOs must use the bandwidth of the existing HFC infrastructure more efficiently to avoid having to upgrade plants as the need for bandwidth increases.

The demand for bandwidth has motivated MSOs to find ways to improve the efficiency of bandwidth utilization. The downstream for DOCSIS and other video services use QAM modulated MPEG-2 transport streams to carry the data. Because of the common transport encapsulation and modulation technique, a single set of MPEG-2 based QAM devices and associated RF bandwidth can potentially be shared across all of these services for more efficient bandwidth utilization. In later section of this paper, quantitative analysis will be given to show the potential saving by sharing QAMs.

However, the current network architecture shown in Figure 1 makes the QAM resource sharing impossible. First, each service is responsible for managing the QAMs dedicated to that service. Thus it is not possible to dynamically share HFC bandwidth between services. Secondly, the DOCSIS CMTS bundles both upstream and downstream together in the single logical device. This makes it difficult to share downstream RF resource between DOCSIS and video services.

A new architecture is now evolving which addresses the problems of the existing architecture. Figure 2 shows the new architecture that is capable of dynamically sharing downstream QAMs between VOD and DOCSIS services. In this new architecture, the downstream QAM is capable of both video and DOCSIS processing. In addition, the function of the DOCSIS CMTS is now broken into 3 separate components. They are the DOCSIS MAC processor, an upstream QAM, and a downstream QAM. These 3 components together are called the modular CMTS (M-CMTS). The components

of the modular CMTS are connected using Gigabit Ethernet. This architecture makes it possible for the downstream QAM to accept both VOD and DOCSIS traffic. In later section of this paper, more details will be given on the data plane and control plane architecture which has the promise of sharing QAM resources.

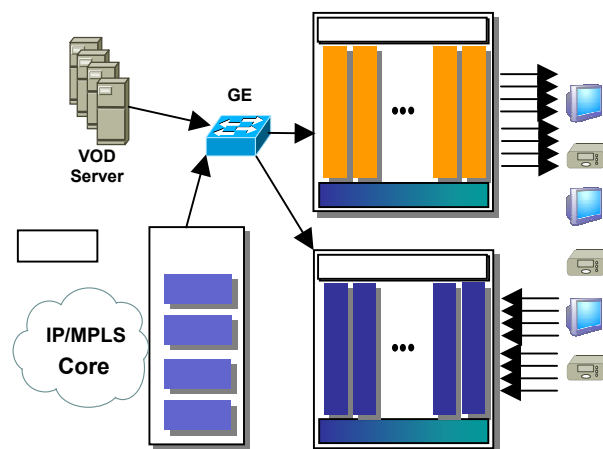


Figure 2. New architecture with universal edge QAM

BENEFIT OF UNIVERSAL EDGE

How much RF bandwidth associated QAMs can be saved when multiple services share the universal edge QAM? This section will do quantitative analysis to show the potential significant savings. We will use switched broadcast, VOD and DOCSIS data services as examples.

The first factor that allows bandwidth savings is the fact that the busy hour associated with each of these services may not be at exactly the same time or day. For example, while the busy hour associated with broadcast and on demand services are typically during the broadcast network prime time period (8:00 PM to 10:00 PM), the busy hour associated with internet access services may occur later in the evening after children have gone to bed.

The second factor that allows bandwidth savings is the savings associated with the semi random behavior of individual subscribers and the savings that can be obtained by taking advantage of the probability distribution of the behavior of a population of subscribers across services. The telephone industry has long used probabilistic models based on subscriber behavior to determine how much bandwidth needs to be deployed to allow a group of subscribers to gain access to the network with a high probability of obtaining network service. A well known model for this type of telecommunication traffic design and analysis is called the Erlang model.

Multi Service Erlang Analysis

From mathematical point of view, Erlang model has provided further evidence that the second factor mentioned above achieves bandwidth savings. The Erlang model shows that the efficiency level of a resource such as telephony trunks or RF bandwidth in a cable plant increases as the number of subscribers that share that resource goes up. When RF resources are shared across multiple services for a given population of users, the effect is to essentially increase the number of subscribers that are sharing that RF bandwidth. This more efficient usage of the RF bandwidth allows less RF spectrum to be allocated to the combined set of services that would be the case if resources were not shared.

While most Erlang calculators are specific to telephony, the calculation itself is applicable to any form of service where the arrival rate of new requests during the period over which the calculation is run (the busy hour) can be assumed to be random. Since the exact timing of how subscribers make VoD requests is not synchronized to external events (such as the advertised beginning of a television show) the arrival rate for VoD requests can be assumed to be random just as

it is with telephony. Another factor that makes the Erlang calculation applicable to video is that the calculation is independent of call hold time. In telephony, the call hold time is the amount of time an average call lasts. It is typically a couple minutes. In VoD, the equivalent of call hold time is the amount of time the typical subscriber spends watching a movie. While this time is likely much longer for video than for telephony, the Erlang calculation is independent of this factor.

Since Erlang B calculators perform calculations in the context of telephony requirement, the variables of a typical Erlang B calculator must be translated into appropriate units relevant to other services such as video and Internet Access. Erlang B calculators are typically used for call center analysis and are readily available from many sources including the Internet. An Erlang B calculator has 3 variables associated with it. The calculator typically allows the user to specify 2 of the variables and it calculates the third variable.

Erlang B Analysis for Video Services

The 3 variables in an Erlang calculator are: busy hour traffic (or Erlangs), blocking factor, and capacity measured in number of lines. Busy hour traffic (BHT) is the number of hours of call traffic during the busiest hour of operation in the system. For VoD services, this can be determined by multiplying the number of homes in a service group, the percentage of homes subscribed to the service, and the engineered peak usage rate for the service. The blocking factor for VoD services specifies the percentage of time that VoD requests will be allowed to fail due to lack of QAM bandwidth. Note that the blocking factor for VoD is usually specified to be very low since it is undesirable to disallow service to a subscriber. The number of lines is the value that we are solving for in the Erlang calculations shown in this paper. For

telephony, this is the number of telephone lines that must be installed to support the specified traffic at the given blocking rate. For video, the number of lines can be translated to the number of video streams that you need QAM bandwidth for. To turn video streams into a bandwidth value, we assume that each video stream requires 3.75 Mbps of bandwidth. To determine the number of QAMs required, we then divide the resulting bandwidth by the bandwidth per QAM (38 Mbps) and round to the next higher integer.

Given the above factors, the full formula for determine the number of QAMs required in a service group for VoD services is:

BStream = BW per VoD Stream = 3.75 Mbps
 BQAM = BW per QAM = 38 Mbps
 Homes = homes per service group
 SR = Subscription Rate
 PR = Peak Usage Rate
 BF = Blocking Factor

$BHT = Homes * SR * PR$

of QAMs =
 $\text{roundup}(\text{ErlangB}(BHT, BF) * BStream / BQAM, 1)$

Note that the above Erlang analysis can also be used for switched broadcast services.

Erlang B Analysis for Internet Access

Erlang B analysis can also be used to model traffic associated with an Internet Access service. While the traffic patterns associated with Internet Access are different than telephony or video, you can still model the Internet Access service as one where subscribers are randomly making requests and the service provider is trying to provide a user experience where the subscriber gets a minimum bandwidth for a certain percentage of the time. The percentage of time that the subscriber does not get this minimum

bandwidth can be considered the blocking factor for the Erlang calculation. The blocking factor for Internet Access can be quite high since the effect of a “blocked” user is that his Internet Access service appears slower than the minimum rate. Another factor that must be taken into account is that when a subscriber is using their Internet connection, they are not always making requests that require bandwidth. We must take this factor into account when calculating BHT. We call this factor the Internet usage factor. The Erlang calculations for Internet Access in this paper will use an Internet usage factor of 20% or .2. Given the above factors, the full formula for determine the number of QAMs required in a service group for Internet Access services is:

IUsage= Internet Usage Factor = .2
 BSub = BW per Internet Subscriber

$BHT = Homes * SR * PR * IUsage$

of QAMs =
 $\text{roundup}(\text{ErlangB}(BHT, BF) * BSub / BQAM, 1)$

Multi Service Erlang B Example

In the following example, we apply the Erlang analysis to a cable plant with service usage patterns that are typical in today’s network. For VoD services, the example shows with 500 homes per service group, a 20% subscription rate, a 10% peak usage rate, a 0.001% blocking factor and peak usage time of 8:00PM.

For the internet access service, we assume 2000 homes per service group, a 30% subscription rate, a 20% peak usage rate, a blocking factor of 1% and a peak usage time of 10:00PM. Finally, to calculate the number of QAMs needed for Internet Access we will assume a minimum rate per subscriber of 1 Mbps.

The non-peak hour usage rate is assumed to be half of the peak usage rate for each service.

Table 1 shows the RF bandwidth requirement and QAM resources needed per 2000 subscribers if these services use statically allocated QAMs. The RF bandwidth calculation must take into account the sum of the peaks of each service. The calculations were done using the Erlang analysis described above.

Table 1. Current RF bandwidth requirement without resource sharing

Service	Usage (%)	Blocking (%)	BHT (hour)	BW (mbps)	QAM
DOCSIS	20	1	24	35	1
VOD	10	0.001	10	101.25	12
Total					13

Note from Table 1 that the amount of QAMs required for VoD and broadcast are much greater than those needed for Internet Access services. Because of this, dynamic resource sharing does not provide much benefit with this type of usage pattern.

Table 2 shows a likely future usage pattern that will become common as the need for Docsis bandwidth grows. The basic assumption here is that the amount of bandwidth that the MSO sells the subscriber for Internet Access service will increase from 1Mbps to 4 Mbps. An example change that will drive the need for higher Docsis bandwidth is the evolution of Web based Video over IP to higher screen resolutions. In this scenario, the video usage is the same as in Table 1, but the following Docsis usage patterns apply. The increased use of Docsis bandwidth will drive down the size of the serving group for Docsis to be identical to that of VoD. In this future example, we assume an Internet Access service with 500 homes per

service group, a 30% subscription rate, a 20% peak usage rate, a blocking factor of 1% and a peak usage time of 10:00PM. From table 2, it is clear that savings can be achieved if the peak usage times for Docsis and VoD services are not the same.

Table 2. Example future RF bandwidth requirement without resource sharing

Service	Usage (%)	Blocking (%)	BHT (hour)	BW (mbps)	QAM
DOCSIS	20	1	6	52	2
VOD	10	0.001	10	101.25	3
Total					5

If QAMs are dynamically allocated between Docsis and VoD, the combined service group can be provisioned for each service peak independently. Dynamic allocation will ensure that the correct number of QAMs is allocated to each service is it reaches its peak usage.

Table 3 and Table 4 shows the bandwidth savings that can be obtained in the using the usage data above if RF bandwidth is dynamically allocated to each service.

Table 3 shows the RF bandwidth requirement at 8:00 PM when the VoD service is running at its peak rate while the DOCSIS service are running at its non-peak hour rate.

Table 3. RF bandwidth requirement at 8:00 PM with dynamic resource allocation

Service	Usage (%)	Blocking (%)	BHT (hour)	BW (mbps)	QAM
DOCSIS	10	1	3	32	1
VOD	10	.0001	10	101.25	3
Total					4

Table 4 shows the RF bandwidth requirement at 10:00 PM when the Docsis service is running at its peak rate while the VoD service is running at its non-peak hour rate.

Table 4. RF bandwidth requirement at 10:00 PM with dynamic resource allocation

Service	Usage (%)	Blocking (%)	BHT (hour)	BW (mbps)	QAM
DOCSIS	20	1	6	52	2
VOD	5	.0001	5	105	2
Total				142.5	4

From the above tables, we can see that dynamic resource sharing between VoD and DOCSIS services requires 4 vs. 5 QAMs to be deployed in the serving group which results in a 20% reduction of plant bandwidth used by these services. The 20% saving comes from the difference in peak hours between the different services.

While not shown in the above example, additional savings can be obtained by sharing the same service group with switched broadcast services. In this case, additional savings can be obtained by using a single QAM pool for both VoD and Switched Broadcast services. These additional savings occur because the number of effective users sharing the same pool of QAM resources is increased. The savings is essentially due to the law of large numbers which is what is represented through Erlang analysis.

DATA AND CONTROL ARCHITECTURE

The clear separation of data plane and control plane components makes it easier to put a common resource manager to manage the resources associated with multiple services. In this section, a data plane architecture that can be used for resource sharing will be discussed first followed by a

control plane architecture that can be used for dynamic resource sharing.

Data plane architecture

There are multiple ways to achieve the resource sharing among different services. Figure 3 shows an example data plane architecture. In this architecture, a VOD server, a real time broadcast encoder, a CMTS core, Downstream QAMs, and Upstream QAMs are all inter-connected through a Gigabit Ethernet network. The GE can switch traffic from any components to any components. With this architecture, DS QAM resources are shared among all three services. In other words, DS QAM is capable of both video processing and DOCSIS data processing.

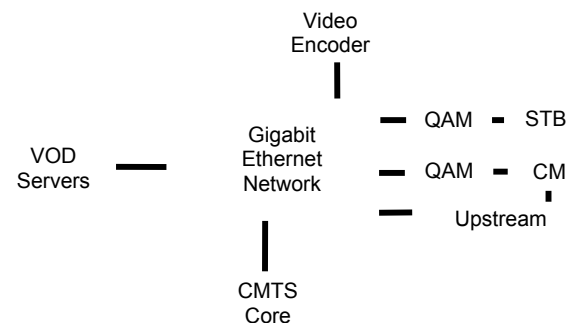


Figure 3. Data plane architecture

This architecture is highly scalable. As more services are added, only the server related to the service needs to be connected to the Gigabit Ethernet. If the QAM bandwidth needs to be increased, additional QAM resources can be shared among all existing services.

Control plane architecture

To achieve resource sharing, a common logical resource management unit needs to exist to coordinate the resource allocation of different services. A component called the session manager is then responsible for

determining the classes of resources required for a session request and communicating with the resource managers responsible for allocating those resources. Figure 4 shows an example control plane architecture that can be used for dynamic QAM allocation. In Figure 4, a control component called edge resource manager is introduced. Edge resource manager is responsible for monitoring the DS QAM resource and allocating QAMs for each service.

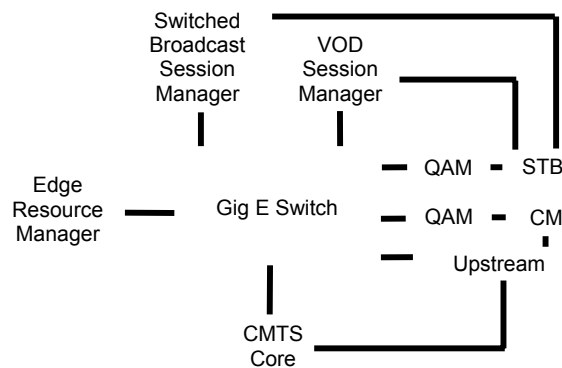


Figure 4. Control plane architecture

The control QAM resources the edge resource manager must communicate with session plane components from each service. In Figure 4, these components are the VOD session manager, and the switched broadcast session manager. The VoD session managers is responsible for accepting user requests from Set Tops for VoD sessions while the Switched Broadcast Session Manager is responsible for accepting channel change requests from Set Tops. Each of these session managers request bandwidth from the edge resource manager as part of the process of instantiating a session.

The CMTS core is responsible for the managing a DOCSIS mac domain. It will request QAM bandwidth from the edge resource manager as part of the process of setting up or modifying the bandwidth associated with a Docsis MAC domain.

The control plane architecture provides two additional functions to the system. The first is QAM discovery while the second is dynamic QAM allocation.

Service discovery protocol allows the Edge Resource Manager to dynamically detect when a QAM comes in or goes out of service. When a new QAM is added or taken out of service, the resource manager will be notified immediately about the resource change. The edge resource manager also maintains a database maintains the mapping of QAMs to service groups.

The second function added is dynamic resource allocation signaling. Each session manager signals to the Edge Resource Manager to allocate or deallocate QAM bandwidth. The Edge Resource Manager returns information of allocated QAMs to each session manager.

This control plane architecture introduces several benefits for the system. First, it simplifies provisioning and management, which in turn reduces the operational expense. In addition, it can improve availability by dynamically reallocating QAMs when a QAM failure is detected. Finally, the separation of session management and resource management make it possible to dynamically allocate QAM bandwidth across services. This provides for more efficient use of existing HFC plant bandwidth.

CONTROL PROTOCOLS

As mentioned in previous section, the control plane supports both service discovery and session signaling. Based on the different requirements for these two functionalities, different control protocols can be selected.

For service discovery, RFC 3219 (TRIP) can be used with minor modifications to suit the needs of cable networks. TRIP is

Telephony Routing over IP protocol which deals the problem of translating telephone numbers into session signaling address of a telephony gateway in VOIP system. When modified for an HFC plant, TRIP allows a QAM to dynamically announce properties about itself to an edge resource manager. These properties include attributes such as the frequency the QAM has been configured for, the HFC service group the QAM is connected to, the amount of bandwidth that is available for the edge resource manager to allocate from, etc.

Dynamic resource signaling could be implemented with a protocol such as RTSP. RTSP is an HTTP based client / server protocol that provide a simple state machine that can be used for resource allocation. RTSP can be used by a session manager to request qam bandwidth from the edge resource manager. The request for bandwidth is encoded in an RTSP Setup message. It includes information such as the amount for bandwidth required for the session / service and the HFC serving group that the bandwidth needs to be allocated from.

After getting the SETUP request, edge resource manager uses its QAM selection algorithm to search for a best QAM to use for this session request. If QAM resources are

available, the resource manager will notify the session manager about the QAM that was selected for the session. If no resources are available to satisfy this session request, the session manager will get an RTSP response with an indication of why the request failed.

CONCLUSION

This paper describes the trend in the cable industry to move to a distributed architecture where RF resources for different services can be shared. From the quantative analysis of this paper, it is clear that resource savings can be achieved by dynamically sharing resources among different services.

This paper further describes a possible architecture to achieve resource sharing and related control plane supports. At the time of writing this paper, the cable industry is actively working on standardizing this architecture and related data / control plane interfaces.

REFERENCES

1. Erlang calculator:
www.erlang.com/calculator
2. Telephony Routing Over IP (TRIP), RFC 3219, IETF, Jan 2002
3. Real Time Streaming Protocol (RTSP), RFC 2326, IETF, April, 1998

CABLE IP VIDEO DISTRIBUTION

Timothy O'Keefe, Robert Sayko and C.J. Liu; Edited by Sean Carolan
AT&T Labs

Abstract

Advances in IP technology are providing Cable Operators with the opportunity to offer innovative interactive services while sharing transmission capacity between IP and traditional video services. This paper will examine how QoS, MPLS, VPN, and Multicast standards for IP networks fit against the demands for quality video distribution, and will explore how these technologies can be used to move nationwide video content distribution to a national IP-based network.

INTRODUCTION

Cable operators are connecting their local networks with IP services to facilitate Internet Access and Voice over IP services. Within their local networks, operators are using Ethernet to distribute VoD streams. The Next Generation Network Architecture (NGNA) is introducing new applications for IP technologies in the delivery of services. NGNA proposes the delivery of a greater volume of video services, and the routing of video service control signaling, over IP based connections.

Advances in IP technology are enabling these new services, and we envision more innovative services and service integration using IP capabilities for the future. Implementations of QoS, MPLS, VPNs and new codec schemes are enabling video services over IP. This paper will examine how QoS, MPLS, VPN, and Multicast standards for IP networks fit against the

demands for quality video distribution, and will explore how these technologies can be used to move nationwide video content distribution to a national IP-based network.

WHY DISTRIBUTE VIDEO OVER IP NETWORKS?

When cable operators offered only traditional video service, satellite access linked their local distribution network access to nationally distributed video content. New Internet Access and VoIP services are fundamentally different than traditional video service; they are more effectively delivered using terrestrially based IP networks, rather than latency-encumbered satellite delivery. Because Internet Access and Voice services send traffic throughout the world, operators need to connect their IP networks to the national and worldwide networks that form the greater Internet. Cable operators are servicing this need by connecting their headends to national terrestrially based IP networks to deliver these services.

For IP networks, scale offers significant cost advantages. A fiber connection into a national network can expand from 100Mbps to 10 Gbps just by upgrading the equipment used at the connection's endpoints. While the bandwidth can increase 100 fold, the cost difference between the 100 Mbps electronics and the 10Gbps electronics is relatively small. A company could effectively reduce the unit cost of their Internet Access and Voice services for a given headend by expanding their existing IP network capacity, and by using the

expanded capacity to receive video services as well.

The fiber medium provides much easier capacity expansion than a satellite-based system. Fiber path construction typically installs a bundle of fibers. As the capacity of an individual fiber fills, lighting another fiber in the bundle can easily activate new capacity. Dense Wave Division Multiplexing (DWDM) also enables each single fiber to multiply its capacity. This easy access to capacity creates the opportunity to apply less compression of the digital video during transmission, thus receiving higher quality video .

Additional efficiencies are gained by using a common technology across all services. When Voice, Video and Internet Access are combined into the IP network, a common set of IP network equipment (routers, switches, etc) can be used across all the services. Technical engineering and operations staffs then have fewer technologies to master. An Operations center can use common IP based network monitoring and management tools. The common use of IP technology across all the services creates efficiencies throughout the business operation.

Video content suppliers are now more likely to have access to IP networks. A well-established national IP network likely passes as close to content producers as it does to content distributors. The content suppliers can input their content to the IP network at any point along the path. Just as the cable operator's IP connection links their Internet Access and VoIP customers to the world, their IP network connection can link them to the world for access to new, interesting content.

REQUIREMENTS FOR A QUALITY IP DISTRIBUTION OF VIDEO CONTENT

Video Delivery Requirements

The requirements for distribution of Video-over-IP are guided by two realities. The first of these are the technical capabilities that must be achieved to present an acceptable video experience. The second are the business drivers that then impose additional technical requirements.

At the present time, formal technical standards that support requirements of real time Video-over-IP, as promoted through a recognized standards organization such as IETF and compliant with MPEG standards for video encoding, are being developed based on needs that are still emerging from the industry. We recognize the IP transport must facilitate transmission of the MPEG stream such that the technical requirements that apply at the endpoint are met. The goal is to guarantee a delivery across the national IP network that will support MPEG requirements for video performance at the endpoint, and will also offer a guarantee of reliability and availability as it relates to uninterrupted service.

The standard measures for video quality include the following:

- Inter-packet jitter
- Packet loss
- Packet arrival order
- Availability

Latency, while critical for real-time interactive applications like telephony, is not as critical for video distribution. Latency of a well run terrestrial IP network will be substantially less than that of satellite television. Inter-packet jitter - the variance in latency from packet to packet – can, for

any reasonable jitter that is expected in an IP network, be accommodated by buffering of IP packets at the receiver. Packet arrival order can also be overcome through buffering and packet reordering in the IP endpoints.

Packet loss is a critical measure of performance for IP video distribution. While transport protocols, like TCP, can request retransmission of lost packets, this is not practical in a video distribution application where a single source is sending a multicast stream of packets to a large number of endpoints. The stream must support transmission, and associated IP overhead, of up to 3.75Mb/s of MPEG2 data for a Standard Definition ATSC transmission; 19.3Mb/s for a High Definition ATSC transmission. A single lost IP packet can translate into a loss of 7 MPEG2 packets.

Network Reliability

Service availability is another key technical capability required to offer broadcast quality content delivery to a cable headend. Since satellite is typically the delivery method for all channels, satellite reliability can be considered a benchmark for video service delivery. Operators don't think about availability on an individual channel basis when all channels arrive via the same transmission path.

However, when various delivery options exist and slight changes in availability performance can be traded against economic benefits, a range of acceptable availability might become part of the delivery system decision. For example, IP networks can be designed for 99.99% availability and 99.999% availability. These two metrics represent a difference of 48 minutes of availability in a year. Yet they represent a

significant difference in cost of network to design, build, and operate. Given the option to make the availability/distribution-price decision on a per channel basis, a cable operator might elect to receive some channels on the lower cost connection and others on the higher cost connection. If satellite is a third, more reliable service delivery, the cable operator might elect to move some lower value channels from satellite deliver to IP network delivery in order to create satellite bandwidth for more high value content.

While no ubiquitous industry standard currently exists for this, it is expected that a video delivery system should employ an architecture that would meet very high level of availability. While a 99.999% availability metric represents 5.3 minutes of outage in a year, cable operators expect interruptions to be few, of short duration and restoration must be seamlessly engineered so that continuity of the video is preserved for the user. Based on these criteria, a successful re-convergence of an interrupted video stream should occur within 1 second or less of detection.

Performance benchmarks and seamless end-user reliability guarantees are not all that drive the IP video distribution technical requirements. In fact there are several other important business criteria that drive addition technical requirements that must be met in order to offer a viable service. These criteria are:

- Costs
- Competition
- Functionality
- Efficiency

The costs of converting to and using Video-over-IP distribution must meet certain thresholds for initial investment (e.g.

CapEx) and the ongoing cost of running the business. The first threshold is met when the reliability and performance of the IP distribution reduces the investment in expensive terminal equipment to control and groom the received video product. The ongoing costs are even more important to the MSO and are also tightly coupled with the three other criteria: competition, functionality and efficiency. Satellite distribution is not the only competition for terrestrially-based Video-over-IP; other competition comes from satellite broadcast providers (e.g. Dish Network, DirectTV). The level of service offered by these competitors is what actually establishes the market's benchmarks for the Video-over-IP distribution service. The bottom line is that IP video distribution must at least marginally beat the cost of alternate providers for equivalent services or provide significantly better service(s) at only incremental cost increases

One key advantage that Video-over-IP must provide over these competitors is the ability to deliver a significantly larger (almost limitless) amount of content. A starting point for the service should begin where the competition leaves off. For example:

- 200 – 400 channels of SD programming
- 25 channels of HD programming
- 100 channels of CD quality music

Functionality is another important consideration for the MSO. Frequently, a Conditional Access connection is required to support the video service delivery, and the CA authority is located away from the headend. These implementations today need to coordinate access to the satellite video services through a separate terrestrial CA network connection. As previously

mentioned, the IP connectivity is two-way. It can already support VoIP and Internet Access service. The same IP connection can be used for the CA connection. The use of IP for video service delivery and CAS may create new opportunities for CA mechanisms.

The final criterion, efficiency, is very important to the MSO as video content options expand, consumers become more sensitive to picture quality, and new innovative services create interaction between video, voice, data, and wireless services. The video service delivery must be capable of expanding overall capacity rapidly, support varying levels of video program compression to ensure high quality content, and provide delivery protocols that easily inter-work with the other services. Finally, it should be easy to add new endpoints to the video service distribution to facilitate easily adding content suppliers and local content distribution networks.

AN OVERVIEW OF IP TECHNOLOGIES

Listed below are the various IP protocols that can be applied to allow video streams on an IP network meet the requirements stated above. These brief descriptions are here just to provide a high level review of the terms, as they will be used throughout the paper. More detailed explanations can be obtained from the IETF and vendor web sites.

IP Encapsulation of Video Frames

MPEG2 frames can be encapsulated in an IP packet. Seven MPEG2 frames are typically combined into one IP packet since this creates a packet size within the 1500 byte limit of Ethernet and enables the packet to be easily moved between layer two transmission protocols. While MPEG2 is

currently the most common video stream protocol, other encoding protocols are also easily placed into IP packets. These IP packets are typically sent as UDP frames. The UDP protocol does not include the ability for the receiving end point to request the retransmission of a lost packet. A national network with 50msec latency might actually allow for 100msec round trip to be used to retransmit a packet. However, Forward Error Correction (FEC) and interleaving packets are more common methods for correcting for lost packets. The Real-Time Protocol (RTP) is also utilized to help sequence packets on the receiving end. If two packets arrive out of order, the sequence numbers in the RTP protocol will allow the receiving end to assemble the video stream in the correct sequence.

Quality of Service (QoS)

QoS standards allow traffic to be marked for specific handling when the network is congested. The most basic handling of IP traffic is called “best effort”. There is no special handling of this class of traffic. The network will try its best to get the packet through as fast as the network will allow. Packets are processed in the order they arrives, first in/first out. The highest quality for traffic handling is called Real-time class. This class gets the top priority from the network equipment. This class of traffic will only be dropped if all the capacity allocated to this class is consumed. Network operators typically allocate enough capacity to this level to support all Real-time traffic they’ve agreed to accept to avoid any dropped packets once it has entered their network. Other classes between Best effort and Real-time define specific behavior during network congestion periods to prioritize and drop packets based on the needs of that traffic class.

Virtual Private Network (VPN)

This segregates traffic on the network such that the network operator can keep some traffic flows separate from each other and from the public Internet. Public Internet traffic just reacts to the IP destination, and will pass any packet to any requested destination. VPN allows the operator to establish additional rules for traffic flow. These rules can include restricting participation in the traffic flows, encryption of the packets in a particular flow, and packet routing based on VPN identification instead of IP address,

Multicast

Most traffic flows in an IP network are point-to-point transmissions. A single source wants to deliver a packet to a single destination, a Unicast flow. Multicast is used when one source wants to send the same information simultaneously to multiple recipients. The multicast routing protocol builds a tree distribution map for all the recipients on the network. A single copy of the packet moves through the network until it arrives at a branch in the distribution tree. At the branch, the network duplicates the packet and sends one copy of the packet down each branch of the tree.

MPLS Traffic Engineering (MPLS-TE)

Traffic Engineering is an extension of the MPLS standard that provides the network operator with more control over the path packets take through the network. This control serves to aid capacity management as well as fast failover recovery. The operator can define specific paths for defined MPLS flows. TE also allows the operator to define a specific failover path for an MPLS flow. This pre-defined failover has become known as Fast ReRoute (FRR).

FRR enables the network to recover a data stream in < 100msec because the alternate path is already known. Since it was predefined, the network doesn't need to take time to discover alternate routes.

POTENTIAL IP IMPLEMENTATIONS FOR IP VIDEO CONTENT DISTRIBUTION

The most basic implementation would be to Unicast an IP video stream across a wide area IP network. This implementation would require the video stream to be replicated at the source for each headend destination. It provides little if any protection of the content, implies a best effort delivery and does not scale very well when there are multiple destinations. It's easy to see that this doesn't fit very well against requirements for quality video programming distribution. The following is a discussion of how other protocols can be applied to improve upon this basic implementation.

Model A: A Multicast solution.

We can make the basic implementation more efficient by implementing multicast. With video services, we expect a single source and many recipients. Multicast provides the ability for the network to take in one video stream and distribute it to multiple recipients by replicating the stream only when necessary. This greatly reduces the capacity demands from the basic implementation and enables the operator to transmit many channels to many end-points within a reasonable network capacity allocation. For example, one allocation of 2 Gigabytes could distribute 300 channels for SDTV and 40 channels of HDTV to any recipient connected to the network.

Model B: A VPN Multicast.

An IP network configured with VPN and Multicast makes a significant improvement over the basic implementation. Implementing the VPN protocol helps protect the video content from being intercepted by unauthorized parties during transmission. The VPN protocol restricts the traffic to specific end delivery points. The network operator controls access to this traffic stream to authorized recipients.

The application of QoS markings to Multicast IP packets is currently not available across all vendors and routers, so the reliable delivery of these streams will usually be dependent on the network operator allocating sufficient bandwidth throughout the network to avoid congestion delay. In addition, a route failure in a traditional IP implementation can cause up to a 10 second outage while the IP network recalculates routes around the failed connection. (An advanced network can provide SLAs that are significantly shorter.) Once the route is reestablished, it can take many more seconds for the video stream end point to resynchronize, re-establish buffers, and return the video stream to a stable flow. This amount of video stream loss far exceeds our requirements for availability.

Model C: MPLS-TE for Fast Network Recovery and QoS

It is desirable to have sub-second restoration capability for video distribution, which can pose problems when required of an IP network. For conventional IP networks, it typically takes five to ten seconds to have traffic rerouted around connectivity failures; either failed links or failed nodes. It's possible to tune nodes on the IP network (using "hello interval", "dead timer", and by leveraging a calculation of

the hold time between two consecutive SPF calculations) to improve IGP convergence time. However, it's still not easy to shave the recovery time down to a level below one second. FRR (FastReroute) in MPLS TE technology enables fail over time of less than 50 ms; an interval that matches the link restoration capabilities of SONET. Fast Reroute is initiated for a Label-Switched Path (LSP) when the feature is enabled for the associated LSP tunnel as a result of a configuration command on the head-end. The head-end router is responsible for informing all routers along the LSP's path that the LSP is requesting protection. The LSP tunnel head-end control module will keep RSVP informed of the status of the Fast Reroute attribute for all active LSPs. When the RSVP module in a Label Switch Router (LSR) [other than tail end] along the LSP's path learns that the LSP should be protected, it will initiate local Fast Reroute protection procedure to protect the LSP against possible failure of the immediate downstream link. Upon link failure, all protected LSPs switch to the backup path. FRR performs the operations to prevent the downstream routers (still along the path in use by the LSP) from tearing down the LSP, if the failure is also detected downstream.

Content delivery services require more significant guarantees for bandwidth rates and for Quality-of-Service (QoS) from the IP network than conventional IP services do. As currently formulated, the leading IETF-endorsed architecture for QoS maintenance of differentiated services, "Diffserv", is strong on simplicity and weak on bandwidth guarantees. In the case of network congestion events, different services would compete for the available link bandwidth. A strict priority queue that includes bandwidth policing for real-time traffic could be enforced, but packet loss and latency still cannot be guaranteed if the rate

of the incoming real time traffic stream is higher than the available bandwidth for the real time traffic. The network operator needs to know how much video traffic will be coming into the network so allocate the necessary bandwidth through the network. DiffServ-aware Traffic Engineering (TE) is a tool for network operators to implement the appropriate bandwidth allocations. DS-TE is meant to enable computing path per class with different bandwidth constraints, and perform admission control over different bandwidth pools. OSPF extensions for DS TE allow advertisement of unreserved TE bandwidth, at each preemption level, for each class type. In DS aware TE tunnels setup time, LSP signaling includes class type as a tunnel parameter, in addition to bandwidth, label, explicit route, affinity, preemption, adaptability and resilience. Class-type aware call admission control will be performed at each LSR during the DS TE tunnel setup. Rate limiting at the head end of the DS TE tunnel can be configured to ensure the traffic into the tunnel does not exceed the provisioned tunnel bandwidth.

Unfortunately, MPLS does not support multicast. MPLS tagging assumes a packet coming into the network can be mapped to one exit point on the network. MPLS can accept a packet coming into the network from one of many possible entry points, assign a tag representing the appropriate exit point, and efficiently direct that packet to the correct single exit point. Multicast wants to do the opposite. Multipoint processing assumes a packet entry at a single exit point should be distributed to many exit points. Relevant IETF working groups are discussing changes to MPLS that could support multicast traffic, but it may be a year or more before those changes begin to appear in network equipment.

As it happens, MPLS-TE no longer has a monopoly on Fast Reroute; standards bodies and the vendor community are working on Fast Reroute on native IP connections. This approach must also be evaluated and compared to MPLS-TE as standards emerge.

Model D: We need a Solution that incorporates QoS, Multicast, Fast Reroute, and VPN.

Possibly the best solution is a combination of all these technologies. We need multicast to make efficient use of network capacity. We need QoS to ensure consistent, on-time delivery of the packet stream. We need VPN to enable access control. Finally, we need Fast Reroute capability to minimize interruptions to the video stream caused by network failure events. Standards bodies and vendors are working to make the whole combination available.

Point to Multipoint Traffic Engineering Label Switched Path (P2MP TE LSP) is currently being proposed in support of the construction of a Point-to-Multipoint (P2MP) backbone network for multicast services. In such a scheme, a P2MP Label Switched Path (LSP) will be set up between an ingress Provider Edge (PE) and multiple egress PEs; the ingress PE would accommodate a multicast source, and the multiple egress PEs would accommodate multicast receivers. Ingress/egress PEs at the edge of the multicast network will handle subsequent multicast routing. The P2MP LSP will be set up with TE constraints and will allow efficient packet replication at various branching points in the network. The proposed P2MP TE LSP would be established by setting up multiple standard P2P TE LSPs. If each P2P sub-LSP is protected by its backup-tunnel, the multicast video traffic can be protected by

the standard FRR TE mechanism, therefore, ensure recovery within 50 ms in case of link/node failure events.

A caveat is: even though there are obvious benefits of deploying TE tunnels in IP network, there are concerns about its scalability and the complexity it adds to network operation. For a facility based ISP that owns the physical links and infrastructure of its IP network, capacity constraint is a relatively minor issue compared to other ISPs which have to purchase or lease capacity from other providers. It is hard to justify sending all IP traffic into fully meshed TE tunnels ubiquitously deployed for a facility based ISP. Instead, only special traffic, such as VoIP, broadcast video, video conferences, or VoD transported in IP network, are candidates to be carried in MPLS TE tunnels. This requires the traffic that enters a configured MPLS TE tunnel get preferential treatment over all other traffic by all routers' queuing and congestion avoiding mechanism along the path. TE queues for configured MPLS TE tunnels in every router the tunnel traverses had been proposed. The proposed TE queues for P2P unicast tunnels can be extended to P2MP multicast tunnels.

IN CONCLUSION

Cable operators have made IP protocols an important part of their network services for Internet Access, VoD, and VoIP. NGNA is creating additional opportunities for IP based services in the network. Advances in IP technologies provide are creating the opportunity to move national broadcast video distribution to IP networks. The application of Multicast and VPN with sufficient bandwidth allocation can provide very reliable video distribution and could be used for some channels today. The

application of Fast Reroute capabilities can bring recovery from link outages to around 100msec and make IP video delivery even more reliable. Mixing national video distribution with Internet Access and VoIP traffic creates economies of scale throughout the business. Putting all service delivery on IP enables new opportunities for delivery to end consumers and innovative service integration. Moving national video content distribution to IP networks could be the next big service breakthrough for Cable operators.

REFERENCES:

1. Extended RSVP-TE for Point-to-Multipoint LSP Tunnels, draft-yasukawa
- mpls-rsvp-p2mp-04.txt by S. Yasukawa, A. Kullberg, and L. Berger
2. Requirements for Point to Multipoint extension to RSVP-TE, draft-yasukawa-mpls-p2mp-requirement-01.txt
3. Performance Analysis of MPLS TE Queues for QoS Routing, Yihan Li, Shivendra Panwar, and C.J. (Charlie) Liu, Proceedings of Applied Telecommunication Symposium/ASTC, 2004 (pp.170-174), April 18-22, Arlington, Virginia.
4. MPLS TE Queue Creation as a Mechanism for QoS Routing, C.J. (Charlie) Liu, Proceedings of Applied Telecommunication Symposium/ASTC March 30-April 3, 2003, Orlando, Florida

CONTENT PROTECTION CONSIDERATIONS FOR DIGITAL CABLE READY PRODUCTS AND SECURE HOME NETWORKING

Brad Hunt¹ and Jim Williams²

¹Sr. VP, Chief Technology Officer, ²VP, TV & Video Systems

Abstract

This technical paper highlights several important content protection considerations for Digital Cable Ready products including secure digital outputs, steps to address the "Analog Hole", secure integrated personal digital recorders and secure home networking.

Adding new features in a secure manner will help maintain the viability of cable television in the competitive and expanding market of digital content distribution. It will also better position the cable industry to launch new innovative programming services that can increase revenue, control churn, and expand the subscriber base.

INTRODUCTION

In the recent past, marketplace solutions for content protection and security were developed by Cable MSOs and other MVPDs through independent negotiations and contractual obligations with content providers and receiver manufacturers. As part of its implementation of Section 304 of the Telecommunications Act of 1996, the FCC issued its Second Report and Order in October 2003 that outlined rules and standards for unidirectional digital cable ready products. Through this ruling the FCC committed to the principle of separation of security functions from the base customer premises equipment to support the retail availability of digital cable set-top boxes that consumers could take with

them when they move. With the separation of the security functions, content owners and MSOs no longer have a direct relationship or voice in the construction of cable receiving products.

Content owners continue to have a vital interest in ensuring that all content distribution platforms are secure not only for existing services but also for future envisioned service offerings. Their views concerning the security aspects of a content acquisition device should be incorporated into the device's technical specifications and the content protection related licensing terms. In addition, the process for approving new protected digital output and secure recording technologies must also include a role for content owners.

As directed by the FCC, the Cable and Consumer Electronics industries have begun working with content owners to define the content protection requirements for next-generation bidirectional "Digital Cable Ready" products. These ongoing discussions regarding the bidirectional framework have provided content owners with an opportunity to express and discuss their views on content protection. As recognized in the Broadcast Flag regulation, the ability of distribution channels to attract high value content is enhanced by due recognition of the security needs of content owners.

This technical paper will highlight several important content protection

considerations related to Digital Cable Ready products.

APPROVAL OF EFFECTIVE DIGITAL CONTENT PROTECTION TECHNOLOGIES

As important partners in enabling content distribution over cable, content owners have a legitimate interest and should have a meaningful role in approving new digital content protection technologies in digital cable products. These include new protected digital output technologies and secure recording methods.

Currently, CableLabs has the authority to approve or disapprove new digital output protection technologies and secure recording methods for digital cable products manufactured under the CableCARD Host Interface License Agreement (CHILA) and the unidirectional DFAST license. It is not clear what functional criteria CableLabs uses to evaluate a digital content protection technology. In fact, the use of a fixed set of functional criteria may be too restrictive in allowing for innovation of new content protection technologies. A more effective manner of analysis and approval should be based on marketplace criteria where content owners' views and actions can lead to approvals based on the marketplace performance of these technologies. At the very least, CableLabs should incorporate a more formal process that seeks and takes into account input and advice from content owners as an integral part of their decision-making process.

SECURE HOME NETWORKING – EVOLUTION FROM COPY PROTECTION TO CONTENT PROTECTION

The nature of customer premises equipment is changing -- evolving from one or more independent receivers with analog video outputs to a suite of networked digital devices that have access to shared resources including tuners, mass storage devices, optical media burners, computers and Internet connections. With continually increasing processing power, Internet connection speed, compression algorithm performance and storage capacity, the customer premises equipment suite is becoming a digital processing, communications, storage, and consumption powerhouse ripe with new content usage possibilities for consumers.

Digital content protection technologies ensure that a particular usage model or cable service offer that is purchased through a conditional access system is honored by downstream devices. Traditional content protection technologies have focused on copy protection. As customer premises equipment evolves into a suite of home networked devices and even to devices beyond the home, the content protection system must also incorporate redistribution control.

The typical usage rights that might be granted in a cable environment include the right to make copies, the right to electronically move content around one's home (e.g. to another TV set) and the right to make a physical copy that can be carried beyond one's home. However, unrestricted redistribution of content beyond one's home would be inconsistent with the licensing rights negotiated from the content owners and could undermine the subscriber-based

business model of the cable television industry.

OPENCABLE™ MUST BE UPDATED TO PROTECT THE SINGLE HOME CABLE ACCOUNT

The current [OpenCable™](#) specifications do not provide the ability to distinguish a single digital cable subscriber with a home network from a group of separate households “sharing” a single cable account using wide area networking. In addition, the content protection afforded by these specifications should provide the ability to signal redistribution control information and manage content usage in accordance with that signaling.

The current [OpenCable™ CableCARD™ Copy Protection System Interface Specification](#) does not provide a means for signaling Redistribution Control. In addition to signaling numeric copy control restrictions of Copy Never and Copy One Generation, this interface specification must have a means to signal Redistribution Control when no numeric copy control restrictions are asserted. For example, this could be the case for programming delivered on the Digital Basic Tier, where a Cable MSO optionally wants to encrypt the service to provide protection against theft of service. In this case, the controlled content would be marked in a manner to signal that there are no numeric constraints on copying within the home or to removable media, but the controlled content must be protected by the host device to restrict redistribution beyond the particular cable subscriber’s home, including over the Internet.

IMPLEMENTATION OF A DIGITAL CONTENT PROTECTION DEVICE KEY REVOCATION SYSTEM

The cable distribution system must provide an end-to-end solution for the delivery and processing of digital content protection System Renewability Messages (SRMs). System Renewability Messages are the common name for the messages that contain digital content protection device key revocation information. Device key revocation provides a content protection technology the means to selectively disable the protected digital output of a compromised device (e.g., a non-compliant device created using a cloned device key) without impacting the general functioning of the device. It is therefore a critical component in managing the effective functioning of a digital content protection technology.

Specifically, the cable system must develop a means for efficiently delivering SRMs from the cable head-end to the digital cable receiver. In addition, both the CHILA and the unidirectional DFAST license must contain explicit obligations for the digital cable receiver to perform digital content protection device key revocation processing when validly received SRMs are presented. Since some digital content protection technologies, like High-bandwidth Digital Content Protection (HDCP), do not store revocation lists, the CHILA and DFAST license must explicitly require real-time processing of SRMs. In the specific case of the 5C Digital Transmission Content Protection (DTCP), the CHILA and DFAST license must require that the device implement “Full Authentication” of the DTCP source function, in order to ensure that full SRM processing is done. These are a few of the requirements for insuring digital cable products incorporate digital content

protection technologies that implement an effective device key revocation processing mechanism.

ADDRESSING THE “ANALOG HOLE”

In the process of delivering protected digital content, the content must be converted into an analog video signal in order to support legacy displays that have only analog video inputs. However these analog video signals can be easily converted back to digital without any obligations to preserve and respect the content’s usage rights information. The protected digital content is said to escape through the “Analog Hole”. The challenge for our industries is to determine the best way to support legacy analog displays without creating an unnatural impediment to the migration to digital.

Several key features of digital cable products are important in addressing the Analog Hole:

- Analog copy control signaling implementation;
- Image constraint on unprotected high definition analog video outputs; and
- Selectable output control capability for new business models.

Each of these features is an important content protection function, and in combination, provides a reasonable approach for addressing the Analog Hole.

ANALOGCOPY CONTROL SIGNALING IMPLEMENTATION

One important component of the solution to the Analog Hole begins with the use of a standardized means for signaling copy control information in the analog video outputs of digital cable receivers. The

application of analog copy control signaling, such as analog Copy Generation Management System (CGMS-A) signaling, has been widely implemented for many years in a number of content protection licenses and specifications. This vertical blanking interval signaling allows the conveyance of usage rights in analog video content. Many digital recorders detect CGMS-A in order to manage unauthorized copying. For example, when the CGMS-A state of “Copy Never” (1,1) is detected in the vertical blanking interval of an analog video signal to be recorded, the digital recording is stopped. In order for this signaling to be deployed effectively, it must be generated correctly in the digital cable set top box.

Both the CHILA and unidirectional DFAST license need explicit obligations for the regeneration and the insertion of vertical blanking interval signals for copy and redistribution control. The MPAA has proposed specific language for explicitly defining CGMS-A, Analog Protection System (APS), and Redistribution Control Information (RCI) signaling in these licenses for all analog video format outputs. In order to ensure full protection, analog vertical blanking interval signaling must also be applied both to upconverted standard definition TV programming that is output as a high definition analog video signal and, likewise, to downconverted high definition TV programming output as a standard definition analog video signal. Finally, analog video outputs should not be permitted absent a standardized means for carrying CGMS-A, APS, or RCI vertical blanking interval signaling. This is currently the case for analog RGB VGA computer monitor outputs.

IMAGE CONSTRAINT OF UNPROTECTED HIGH DEFINITION ANALOG VIDEO OUTPUTS

Content owners are very concerned about the introduction of digital recorders that exploit the high definition Analog Hole. The price of high definition analog-to-digital video converter devices is falling and could soon lead to the introduction of consumer devices that digitize and record unprotected analog high definition video content. The use of image constraint on unprotected HD analog video outputs is an important tool in addressing the high definition Analog Hole. The optional use of image constraint on unprotected analog high definition video outputs has not been demonstrated to have any visual impact on legacy HDTV displays having only analog video inputs.

The use of image constraint provides incentives for consumers to use the higher-quality, protected digital interconnects that are becoming available in the marketplace. Since the obligation to implement image constraint is in the DFAST license, all unidirectional CableCARD-equipped host devices being introduced today have image resolution constraint capability. This capability must be implemented in future digital cable products.

SELECTABLE OUTPUT CONTROL CAPABILITY FOR NEW BUSINESS MODELS

Under the unidirectional regulation, the FCC acknowledged that selectable output control could be appropriate for use in the future. Cable is afforded two key benefits by deploying selectable output control capability in Plug and Play products.

First, as suggested by the FCC, selectable output control might enable future

applications that are advantageous to consumers, such as new early-window business models. For example, in order to create a more secure environment for an early-window high definition video programming service, an MVPD may find it advantageous to deliver this service with the requirement that unprotected analog high definition video outputs are disabled and only digital outputs protected with HDCP and DTCP are allowed.

Second, selectable output control could also help address unknown problems, such as patent claims and court orders involving a previously-approved content protection technology.

In order to make these future permitted uses possible, manufacturers should be required to incorporate selectable output control capability in all digital cable products.

CONTENT PROTECTION OBLIGATIONS FOR HARD DISK DRIVE INTEGRATED RECORDERS

Integrated Personal Digital Recorders (PDR) in Digital Cable receivers provide many attractive benefits to consumers, such as pause, time-shifting, and the movement of temporarily stored recordings of “Copy One Generation” programming to removable media. But in order for integrated recorders to provide this functionality, the content and associated usage rights information must be securely and persistently protected and content usage must be effectively managed in accordance with those associated usage rights.

In the case of temporary recordings of “Copy Never” programming, the content must be cryptographically bound to the receiving device doing the recording so that

it is not removable and not itself subject to further copying before it is rendered unusable. The temporary copy should be encrypted in a manner that provides no less security than that of the Advanced Encryption Standard (AES) using 128-bit keys. Since rights associated with "Copy Never" content preclude making a permanent copy, the default expiration time of temporary recordings of "Copy Never" content should be 90 minutes. This also requires that the cable system provide a secure source of time to the digital cable receiver/recorder in order for it to securely manage time expiration of bound copies.

In the case of recordings of "Copy Once" programming by integrated PDRs, many of the same requirements for "Copy Never" content are also needed. In addition, these recordings must be remarked to "Copy No More" to prevent further copies from being made by downstream recording devices.

Finally, one of the most important missing features of current Digital Cable Ready products is the provision of a secure time source and a standardized means for signaling time expiration of bound copies. Incorporating this functionality into next-generation digital cable receivers with integrated recording capabilities is critical in supporting a wider range of time-shift, rental, and sell-through programming options for consumers.

LABELING STANDARDS FOR UNIDIRECTIONAL AND BIDIRECTIONAL DIGITAL CABLE PRODUCTS

Based on the bilateral-negotiated DFAST license, a broad array of unidirectional Digital Cable Ready products are beginning to be sold in the marketplace. Even though these devices incorporate a CableCARD

slot, they will not be able to access interactive programming services, such as interactive Video-On-Demand (VOD) and impulse Pay-Per-View (PPV) offerings. If a successful conclusion is reached in the cross-industry bidirectional digital cable negotiations, a new bidirectional framework will be created producing a new generation of bidirectional digital cable ready products that incorporate advanced content protection, copy management, and device programmability. These features will better enable cable operators to provide a wide range of new interactive programming services, including early-window content, to cable subscribers purchasing these new bidirectional devices.

However, content owners are concerned that consumers must be properly educated about the more limited set of programming services available to a unidirectional digital cable receiver as compared to the wider range of new, interactive services that will be available to subscribers purchasing bidirectional digital cable products. Although the current market availability of unidirectional devices is helping to facilitate the Digital Television transition, content owners believe that consumer electronics manufactures and consumer electronics retailers must accept the responsibility for clearly labeling digital cable ready products and for educating consumers about the programming and interactive service availability differences. This is critical to help the customer make an informed purchase decision when considering whether to buy a unidirectional or an advanced bidirectional digital cable ready product.

SUMMARY

This technical paper has highlighted several important content protection considerations related to Digital Cable

Ready products. Addressing these issues is a critical step in maintaining the viability of cable television in the competitive and expanding market of digital content distribution. It will also better position the cable industry to launch new innovative programming services that can increase revenue, control churn, and expand the

subscriber base. Content owners look forward to continued collaboration with the cable and the consumer electronics industries in addressing these issues that will lead to the introduction of exciting new digital cable products and program service offerings for consumers.

CONVERGED DATA NETWORK ARCHITECTURE (CDNA)

Michael J. Emmendorfer
Charter Communications

Abstract

What is CDNA? First, CDNA is the migration of all services (Voice, Video, Data) both serving residential and commercial customers over a Common Network Architecture, including optical transport, Internet Protocol (IP) core, IP distribution, various IP/media access layer technologies (coax, fiber, or wireless). Second, CDNA is a methodology with the objective to reduce layers of the network by converging component functionality within the network elements, as well as increasing the remaining network functionality to support any service over any access layer technology. Third, CDNA is the convergence to a Virtual National Backbone using MPLS (Multi-Protocol Label Switching) technology.

In summary, CDNA takes a holistic view of the network and has three fundamental principles: 1.) Convergence of Services over Common Network Architecture, 2.) Convergence of Network Layers (Reducing Network Elements and Increasing Service Functionality), 3.) Convergence to a Virtual Core Backbone using MPLS Technology

This white paper will illustrate an overview to our Converged Data Network Architecture, including our evolutionary plans, and benefits as well as challenges to this target network architecture.

CONVERGED DATA NETWORK ARCHITECTURE (CDNA)

Executive Summary

Charter Communications – like most Multiple System Operators (MSOs), Regional Bell Operating Companies (RBOCs), and Service Providers – is preparing for an “All IP World,” and is at various stages of implementing this IP strategy throughout our networks. This white paper is an overview of Charter Communications’ strategy and approach for a Converged Data Network Architecture (CDNA). The network will need to support all voice, video, data products serving residential and commercial customers over one common infrastructure. Some of these service and lines of business include Internet Data Services, Telephony, Interactive TV, Video On Demand (VOD), All Digital Video, Multi-Media Services, and a host of local and national based Commercial Products and Services.

The CDNA target architecture will continue to evolve over time. It is important to note that CDNA is not limited to coaxial-based technologies and architectures as stated in the abstract, but rather a holistic view of the network end-to-end with target architecture to support any service through any network technology (coax, fiber, or wireless).

The Convergence of Services over common network architecture, such as commercial and data services across an MPLS enabled CMTS, as well as

switch/router platforms, will allow for the replacement of historic services such as Frame Relay, and has been deployed since 2002. Charter views MPLS first as a service or revenue enabling technology, second as a network label switching technology, and finally the ability to logically partition services across the network. We have continued to converge some services over our IP distribution and optical transport network infrastructure, which includes Residential Voice and Data Services, VOD, Simultrans (All Digital Video), and Commercial Services to various enterprise customers

The Convergence of Network Layers in the metro area will reduce the cost to the MSO, and our target is to enable access layer elements as well as distribution layers with DWDM optical interfaces, by-passing the traditional transport layer as well as aggregation routers. The third generation cable platform providing video and data service is being formulated. There are new revenue generating services and the ability to create CAPEX and OPEX expenditures by enabling the legacy access layer platforms (CMTS) to support a full range of commercial and residential products lines.

Charter's vision is to create convergence to a Virtual Core Backbone implementing MPLS Technology, and will allow us to take the lead in various data, telephony and commercial services.

The opportunities and benefits of CDNA have just begun to emerge, and others will be realized over time. We are able to leverage multiple Lines of Business (voice, video, and data) to support the network infrastructure, and improve the economics for all lines of business, both from a capital and operation perspective.

The benefits to a CDNA strategy will create the vehicle for integration of services, improve the customer experience, increase loyalty, and defend our customer base. The CDNA architecture is a migration away from proprietary systems to standards-based products that will enable us to accelerate the advanced service deployment while driving down operational efficiencies in terms of leveraging a common work force and fewer network systems.

There are challenges that we have encountered and options that need to be considered as we make our way to a truly converged network. A few of the challenges include driving full service features on the 1st and 2nd Generation CMTS, defining the 3rd generation CMTS/Edge QAM to support the full breath of commercial services to address competitive threats (a possible PON replacement technology), integration of long reach DWDM optics on distribution as well as the access layer elements with fault management capabilities.

Introduction

This white paper is an overview of a target network architecture based on business, operations, and technical requirements to support Residential and Commercial Data Services, Telephone Services (VoIP), VOD, iTV, Simultrans, All Digital, and other IP-based Products & Services, internally called Converged Data Network Architecture (CDNA).

Charter Communications has convergence at various portions of the network today, and others will emerge over time. This paper has three fundamental principles for a CDNA strategy.

First, we will begin with the Convergence of Services over a Common Network

Architecture supporting Commercial and Residential customers with several product offerings including Data, Voice, and Video. An advantage of convergence of services across a common network is economies of scale and the ability to offer more features and integrated services for the consumer.

Second, the methodology of Convergence of Network Layers with the convergence to IP will inherently force the migration of services from legacy platforms (often proprietary) to IP network elements (which will be standards based). In addition, this second principal of the convergence of the network is increasing functionality and allowing the service provider to reduce layers of the network and also enhance the remaining layers functionality to support the convergence of services across a common network. This can fundamentally change the network architecture.

Finally, the third principle is the Convergence to a Virtual Core Backbone using MPLS VPN Technology. MPLS technologies have been in operation and offered as a service by backbone service providers for years. However, the adoption of this technology in the cable industry has, until recently, been only in the hands of a few across the world; and of those few most are only using MPLS as a network label switching technology and for traffic engineering across a backbone or MSO core network. In 2002, Charter embraced MPLS as a Service / Revenue enabling technology to position services against the RBOCs and Service Providers utilizing MPLS VPN (RFC 2547 bis) as well as other MPLS based technologies.

BUSINESS REQUIREMENTS AND SERVICE DRIVERS

Charter Communications offers four core services that are starting to cross-pollinate with each other as outlined below. This hybrid approach will require more feature sets as the services evolve, but it can be classified into the following service categories: Video, Voice, Data, and Network.

Video Based Transaction Services

- VOD, SVOD, HDTV, etc.
- Migration to a All Digital Network
- Interactive Set-Top and DVR
- Migration to End to End IP Based VOD and Content Services

Voice Based Transaction Services

- Primary Line Voice Services
- Business Class Telephone Services
- Multimedia Voice Service (SIP)
- Integration with Video Products

Data Based Transaction Services:

- Unified Mail Services
- Web Hosting and DNS
- Centralized Storage and Data Back-Up
- PC-based Virus/Spam Protection

Network Based Services

- Bandwidth Speeds (Tiers)
- Quality of Service
- Network Based Virus Protection
- Security and VPN Services
- Bandwidth Management
- Single or Multi-Site Connectivity Products
- LAN Extension Services
- Transparent LAN Services
- Frame Relay Replacement
- Routing Services
- TDM and SONET Services
- Network Storage Transport Services

ENGINEERING AND OPERATION REQUIREMENTS

The CDNA architecture requires centralized management (NMS & EMS) and full Fault Management, Configuration, Accounting, Performance, and Security (FCAPS). A key driver for CDNA is reducing the complexity of the network, reducing the layers of the network, and differentiated products for service delivery. Also, accelerating the convergence of services and network architectures with the requirement of high reliability and availability across standards-based platforms will advance services deployment, and drive operating efficiency.

FUNCTIONAL TRACKS

Charter has taken an approach to partition the end-to-end *service delivery network* into Functional Tracks that represent common types of services and/or technological areas of concentrations. This paper will concentrate on Track 2 – Converged Data Network Architecture (CDNA) with an example illustrated below in figure 1.

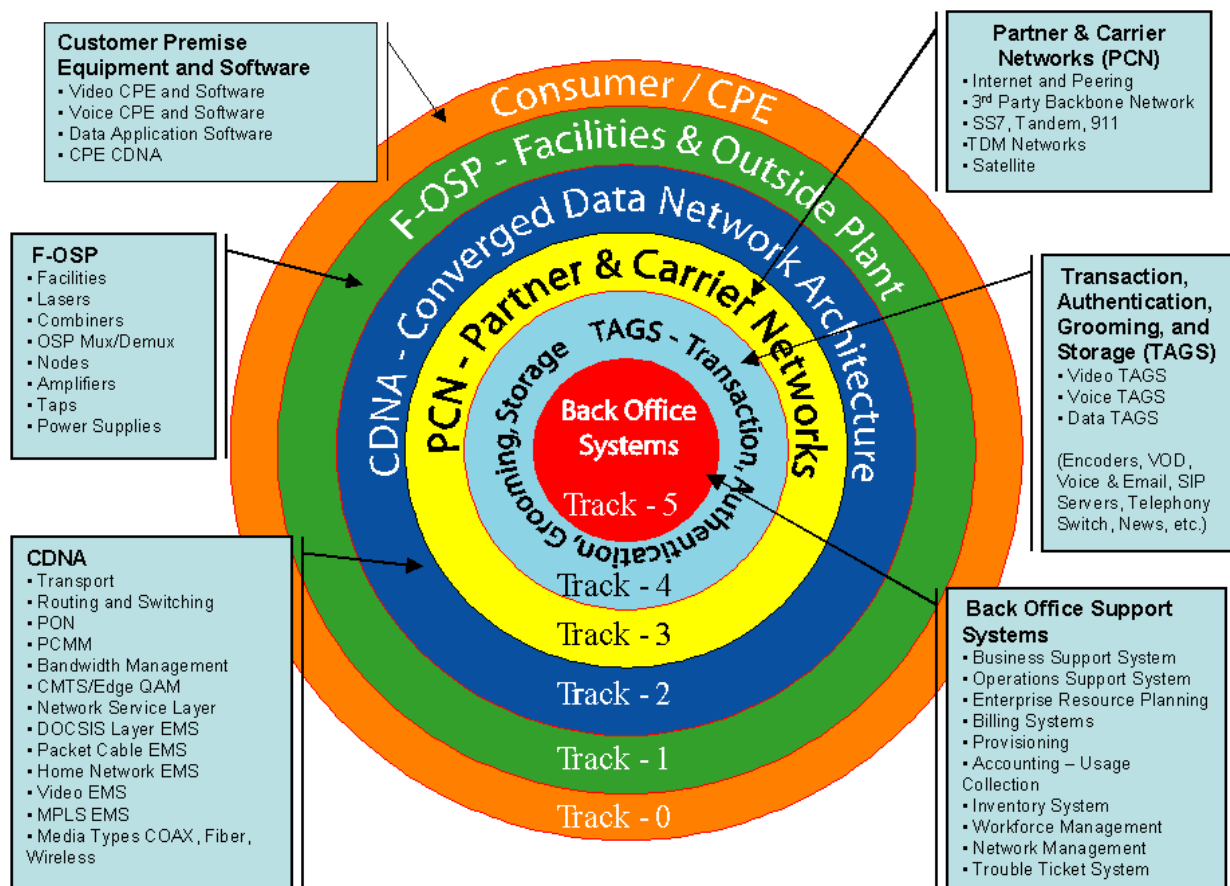


Figure 1: Functional Technical Tracks

THE CDNA TECHNICAL AND MIGRATION STRATEGY

Overview

Reaching our target architecture of full convergence of all services across one network will certainly “not happen over night” since this is an evolutionary path. The following section is a high-level strategy for CDNA, and a migration plan to the target architecture. The target architecture has four layers, including optical transport, Internet Protocol (IP) core, IP distribution, and various IP/media access layer technologies (coax, fiber, or wireless). This next section will examine the Access Layer to increase functionality to offer new services as well as technology improvement to allow seamless integrations. The Figure below is the end-to-end a high-level strategy for CDNA:

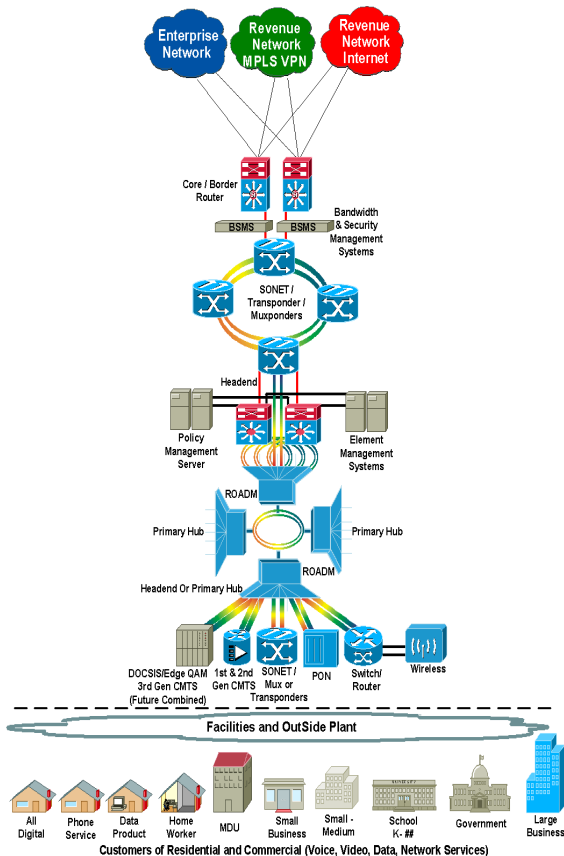


Figure 2: CDNA End-to-End

The figure below is the CDNA access layer component support residential and commercial voice, video, and data service, both end-to-end IP and MPEG.

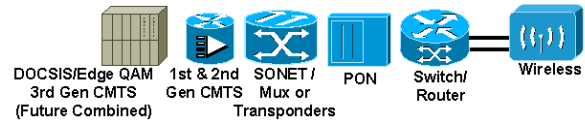


Figure 3: Next Generation Access Layer Elements (Headend and Primary Hub)

Stage 1 Service Enabling Legacy Access

Charter’s objective is to increase functionality of the current CMTS infrastructure with software enhancements, including Sub-interfaces, VRFs, Layer 3 MPLS, MB-BGP, IS-IS Routing Protocol, and others. This will enable the MSOs to leverage an existing CMTS infrastructure for new revenue streams with higher margins. This asset continues with its financial depreciation schedule, but with software we enable a new revenue stream like those to support commercial services; specifically services to replace the incumbent provider. Figure 4 below represents MPLS enabled CMTS to provide frame relay replacement services to support our Commercial Business Unit.

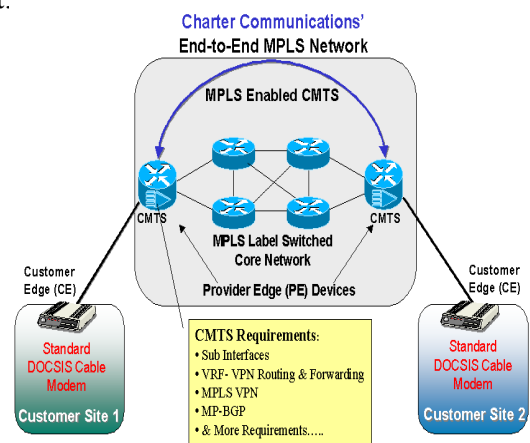


Figure 4: Delivering Commercial Services with MPLS VPN enabled CMTS Delivering

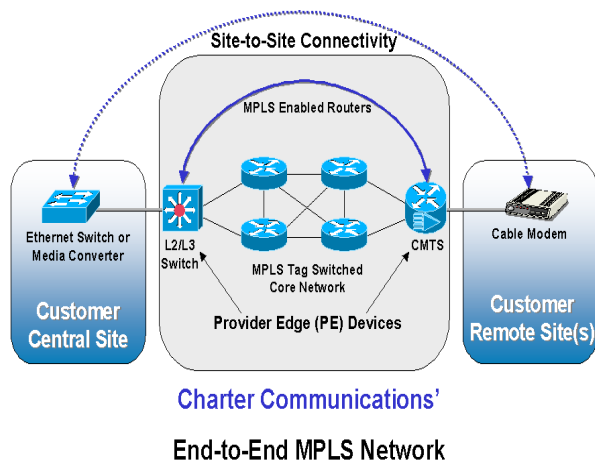


Figure 5: Commercial Services using MPLS VPN enabled CMTS and Switch/Router

The illustration in **Figure 5** enables the customer to use our IP VPN to connect remote site(s) to a central office. This example has the customer provisioning their PCs and SIP phones with their own IP address space, their SIP phones use our provided private network to use 4 digit dialing as well as off-net dialing.

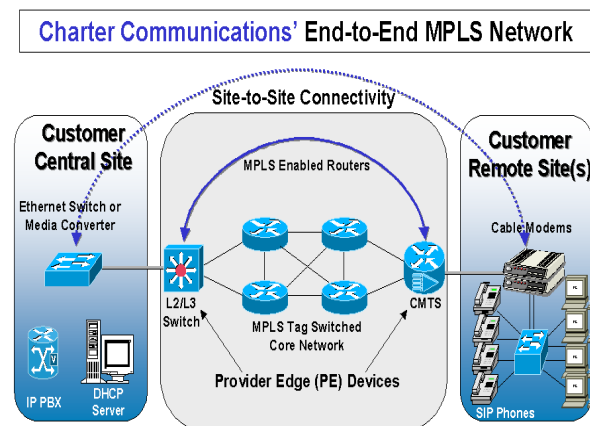


Figure 6: Customer Traffic Transparency across Charter's Network

We are exploring additional service enhancement to the CMTS, which may include Layer 2 Services for Transparent LAN Services, using EoMPLS, QinQ (*802.1q Tunneling), as well as VPLS.

These technologies are traditionally found on switching platforms. There are critical service and management features for the next generations / 3rd CMTS/Edge QAM.

We launched Passive Optical Network (PON) technology in 2002 and SONET services in 2000 to support commercial product offerings such as Ethernet and TDM services, our service enabling strategy is to place CWDM optics in the access network between the Charter facility and the customer, to conserve fiber assets.

Stage 2 Converged IP Distribution for All Services

Charter has placed over its Distribution layer network Residential Data Services, Commercial Services, Voice Service, VOD Services, All Digital video product offering.

Stage 3 Transport Network Migrations

Charter CDNA Architecture will greatly change the Optical Transport network in the metro market; to that end Charter has begun a migration of IP traffic off of SONET as well as RPR over SONET, to transponder and muxponder.

Stage 4 Bandwidth and Security Management System (BSMS)

What is BSMS? BSMS is unified solution for bandwidth management (high bit rate application management) as well as Security Management (Intrusion Prevention Systems, Protection from DDOS, and targeted services security protection). Charter implemented BSMS in our large markets; as a result of this deployment we captured the following statistics:

- Avg. of 40,000 malicious packets for every 100k subscribers per day.
- In one of Charter's larger markets there were 14 million malicious packets events in 30 hours
- Attacker Ratio: For every 1 attack that comes into a Charter Network there are 11 attacks from our subscribers going out

A case study: BSMS Value Assessment and Network Worms

Impact of Network Worms and Malicious Attacks Measurement from July 2003 – November 2003

- **August 11:** The Blaster Network Worm was introduced
- **August 22:** The SoBig Virus Hit was introduced

Table 1: Call Center and Operation Impact from Network Worm (Source: Jon Mandani, Charter Communications)

MONTH	HSD Repair Calls % Increase from July	HSD Truck Rolls % Increase from July
July	-	-
August	70%	26%
September	105%	18%
October	47%	16%
November	8%	(-8%)

We deploy BSMS at the border of our network at the Internet drain; this does not protect or manage traffic to and from our subscribers, which would be on net. As we approach convergence of Voice and Video on one IP network this feature has a considerable increase in importance. Our BSMS strategy is a requirement for the next generations / 3rd CMTS/Edge QAM, Passive Optical Network, and Switch/Router.

Stage 5 Convergence to a Virtual Core Backbone using MPLS

The creation of a Virtual National Backbone using MPLS creates cost savings and revenue generation opportunities thought not possible because of geographic separation from our facilities (Headend, Hubs, and Offices). Charter plans to create a Virtual Backbone using MPLS, using two technical approaches representing two business needs:

1. **MPLS** - Charter becomes a customer of a MPLS-enabled service provider establishing connectivity to remote offices and facilities around the country to support the internal enterprise network. This effort replaces the frame relay technology that is currently implemented.
2. **Hierarchical VPNs** – For Charter to become an MPLS VPN Service Provider would require connectivity across an MPLS VPN enabled service provider nationwide. This enables some key features:
 - a. Commercial Services to provide connectivity and QoS to our customers between our current network footprint as well as outside service areas. Global customers are now possible across this virtual network.
 - b. MPLS enabled service (VoIP and others) to provide connectivity and QoS across carrier networks to other Charter sites or partners.

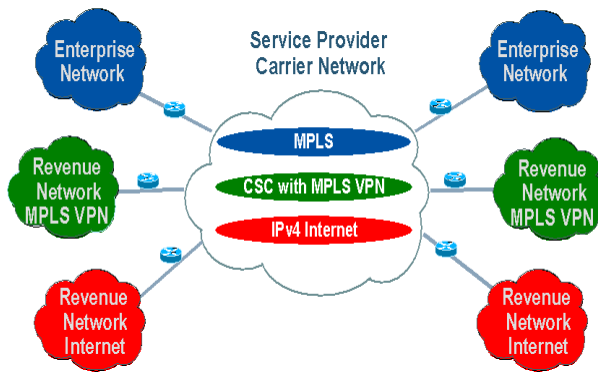


Figure 7: Connectivity Across MPLS Carrier

This opportunity is realized by the ability to offer services where an MSO's backbone or metro network cannot reach without a significant capital investment.

We are considering placing other services across our MPLS Virtual Backbone as well as the Hierarchical VPNs. The diagram on this page illustrates the Internal Enterprise Network (MPLS), Revenue Network MPLS VPN - Hierarchical VPNs, and Internet.

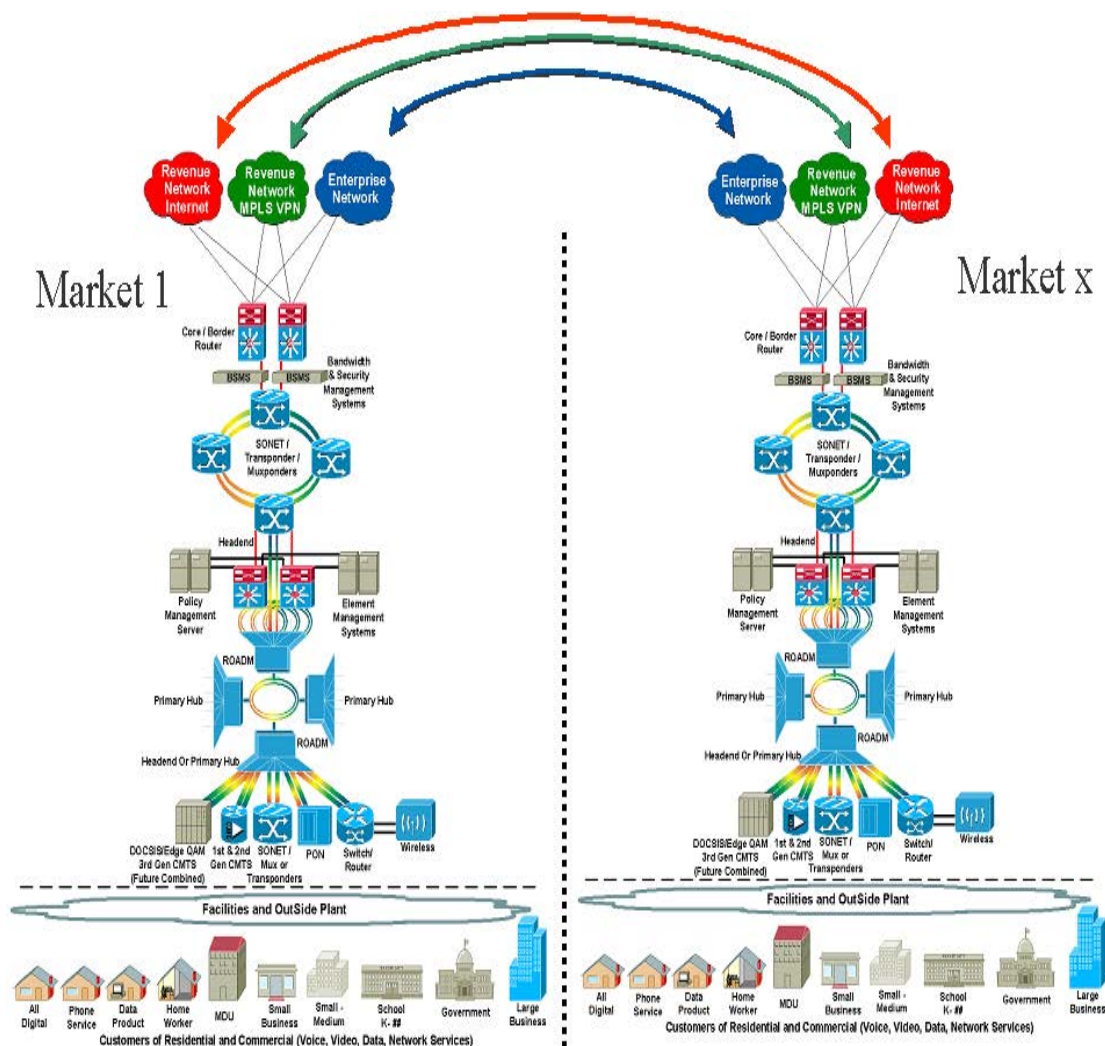


Figure 8: Illustration of Interconnect of two (2) MSO markets connected over MPLS enabled Carrier(s) utilizing MPLS, Hierarchical VPNs, and Internet Access.

There are a multitude of technical considerations for Hierarchical VPNs and MPLS:

- Using the RFC 3107 - Carrying Label Information in BGP-4 to interface with the carrier, aka BGP send label
- DSCP for IPv4 to Carrier MPLS EXP
- EXP MPLS to Carrier MPLS EXP
- Depth of label stacking for your MPLS VPN Carrier
- Multicast over MPLS VPN
- MPLS Monitoring Visibility

Stage 6 ITU DWDM Enable Distribution and Access Layer with ROADM Technology

The integration of long reach ITU Dense Wavelength-Division Multiplexing (DWDM) Gigabit Interface Converter (GBIC) (or SFP) “natively” on access layer elements, like that of CMTS, L2/L3 Switches, PON, Edge QAM, as well as the switch/router distribution layer elements could bypass the transport layer in the metro markets. In fact, aggregation L2/L3 switch/router(s) located at primary hubs would not be required as well, with native DWDM optics on the access layer elements.

Assets already deployed in primary hub locations could be re-allocated from Switch Aggregation and Optical Metro Transport Distribution to Access Layer devices used for revenue generation. **Figure 8** illustrates this proposed design; this architecture reflects the optical transport layer as well as L2/L3 Aggregation Layer Switch at the Primary Hub (s) as these components are not required.

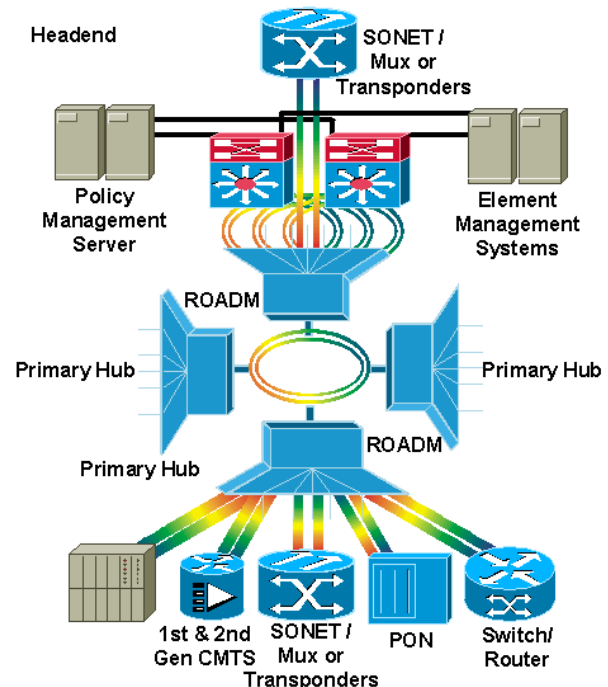


Figure 9: ITU DWDM GBIC (or SFP) “natively” on Access Layer & Distribution Layer

There are a host of challenges with native ITU DWDM Ethernet Interfaces (1) Gigabit and/or (10) Gigabit that will need to be addressed and solutions developed, however once realized the capital and operational savings are significant.

Though work to improve router switch over and convergence time, like those found in RFC 3623 - Graceful OSPF Restart and other similar standards to improving the protection capabilities of the router to be at parity with SONET scheme. However the photonic level impairment measurements protection schemes inherent with SONET/SDH and ITU G.709 are not available, such as:

1. Performance monitoring Fault management, errors, alarms, and performance monitoring like which exist in SONET/SDH and ITU G.709, for optical layer problems.

2. SONET/SDH and ITU G.709 defined threshold are met for link impairments a message for automatic protection switching (APS) occurs.
3. POS interface contains this protection.

Below, is a set of questions to the industry surrounding protection schemes for Ethernet without wrappers (G.709) so that protection scheme like those defined by SONET/SDH and ITU G.709 could be available for native Ethernet?

1. How can ITU DWDM Ethernet Interfaces (1) Gigabit and (10) Gigabit obtain the optical measurement for Errors, Alarms, and Performance Monitoring similar to that of SONET/SDH and ITU G.709, WITHOUT placing a wrapper around the Ethernet frames?
2. How can Ethernet get the equivalent of signal degrade bit error rate measurements (link up but degraded) and issue the equivalent of SONET APS to a Router's IGP to force a re-convergence of a link that is not down and the IGP hellos and dead timers (even if default setting are reduced) does not declare the link down?
3. Keeping Ethernet enacted could an equivalent signal degrade BER measurement over IP link monitor and measure a consistent stream for degraded link and issue an alarm and/or issue a signal to the IGP to converge.

These performance measurements and protections scheme are important with UDP stream (VoIP or Video), TCP data traffic may not detect the impairment, voice and video could be affected.

Reconfigurable Optical Add/Drop Multiplexer (ROADM) is being considered as part of the CDNA strategy for per lambda optical wavelength managements.

Stage 7 The 3rd Generation CMTS Platform

Charter has defined an Edge Layer that can support services and technology, independent of media (coax, fiber, and wireless), this integration of the services will provide greater serviceability of commercial services, arguably one of our industry's fast growing segments. With regards to the next generation coaxial platform, known as the third generation CMTS, also known as the modular CMTS is certainly the "buzz" these days.

As discussed earlier in the paper we use the 1st and 2nd generation CMTS platforms to deliver data and voice service to residential customers, in addition we also offer advanced services to commercial customers. Charter partitions the CMTS logically to create management sub-interfaces for provisioning and management, as well as for revenue services. This partitioning enables Charter to apply routing rules, security rules, or service for revenue generation per logical interface.

Charter is interested in this next generation platform because we see this not just as a platform that can provide integrated Data and Video services, but this has the potential to provide high bit rate services to commercial subscribers and to augment a Passive Optical Network.

BENEFITS

We hope CDNA will provide capital expense savings through economies of scale and standards platform architecture. We believe that operational expense savings and increase network availability, manageability

(remote), and a shared work force supporting many services, but one network should improve the economics of all line of business.

Convergence will enable fewer facilities for complex and expensive transaction processes (Encoding, Storage, Telephony Switching, Email, Hosting and eventually few Headends.

CHALLENGES

There are two core challenges as part of the CDNA strategy, Ethernet Optical management and monitoring, like those found in SONET and ITU G.709. In addition, the feature and functionality of the

3rd generation CMTS/Edge QAM, especially in the areas of logical interfaces for management and security as well as commercial services capability (bandwidth and services).

CONCLUSIONS

Convergence to IP end to end will emerge over time for the cable operator, and legacy technology are increasing using IP for transport services, as this evolves over time we are well positioned to support this migration. This architecture can have significant cost savings.

Michael J. Emmendorfer
Charter Communications
MEmmendorfer@chartercom.com

DOCSIS PERFORMANCE ISSUES

Jim Martin

Department of Computer Science, Clemson University

Abstract

We have developed a model of DOCSIS using the ‘ns’ simulation package. We identify a set of possible DOCSIS performance issues which includes complex interactions between downstream TCP connections and upstream MAC operation, vulnerabilities caused by MAC level denial-of-service attacks and fairness issues. We summarize our ideas involving bandwidth management to address the issues.

INTRODUCTION

The Data over Cable (DOCSIS) Service Interface Specification defines the Media Access Control (MAC) layer as well as the physical communications layer that is used in the majority of hybrid fiber coaxial cable networks that offer data services [1]. A Cable Modem Termination System (CMTS) interfaces with hundreds or possibly thousands of Cable Modem’s (CMs). The original DOCSIS MAC interface (version 1.0) provides a best effort service with simple prioritization capabilities. DOCSIS 1.1, which is currently being deployed, adds a set of ATM-like services along with the necessary QoS mechanisms. The follow on standard, version 2.0, enhances the physical layer communication methods with higher upstream data rates and improved tolerance to bursts of noise.

The CMTS makes upstream CM bandwidth allocations based on CM requests and QoS policy requirements. The upstream channel is divided into ‘minislots’ (referred to as *slots*) which, depending on system configuration, contain between 8 to 32 bytes of data. The CMTS periodically sends a ‘MAP’ message to all CMs on a downstream channel that indi-

cates upstream bandwidth allocation over the next ‘MAP time’. The MAP provides slot assignments for particular CMs in the form of data grants, provides opportunities for CMs to request upstream bandwidth using a contention-based request process and identifies which slots are to be used for system overhead.

A critical component of the DOCSIS MAC layer is the upstream bandwidth allocation algorithm. The DOCSIS specification purposely does not specify these algorithms so that vendors can develop their own solutions. However, all upstream bandwidth management algorithm will share a set of basic system parameters such as the amount of time in the future that the scheduler considers when making allocation decisions (we refer to this parameter as the *MAP_TIME*), the amount of upstream bandwidth allocated for contention-based bandwidth requests and the range of collision backoff times. These parameters are crucial for ensuring good performance at high load levels.

We have developed a model of the DOCSIS MAC and physical layer using the ‘ns’ simulation package [2]. In previous work we reported on the impact of several DOCSIS operating parameters on TCP/IP performance [3]. In this paper we extend those results by looking in greater detail at the impact that the MAC layer has on TCP performance when using the DOCSIS best effort service. We show that the interaction between DOCSIS and TCP exposes a possible denial-of-service vulnerability. By exploiting the inefficient, contention-based bandwidth request mechanism, a hacker can severely impact network performance. We demonstrate fairness issues involving TCP and video streaming protocols that are ‘TCP-friendly’. Most streaming

video applications do not respond to network congestion. The Internet community has addressed this by developing the Datagram Congestion Control Protocol (DCCP) which provides an unreliable datagram transport service that includes TCP-compatible congestion control algorithm referred to as the TCP Friendly Rate Control (TFRC) protocol. While DOCSIS impacts downstream TCP performance, it does not impact the performance of TFRC (at least to the same degree). This causes TFRC flows to steal bandwidth from similarly configured TCP connections. We summarize our ideas on how bandwidth management can address these issues. We propose a bandwidth management algorithm that addresses fairness issues that include controlling TCP unfriendly flows and also subscribers that consume a disproportionate amount of bandwidth.

This paper is organized as follows. The next section presents the operation and features of our DOCSIS model. We present experimental results illustrating the performance issues. We then present our bandwidth management algorithm. We end the paper with a discussion of related work, present conclusions and identify future work.

SUMMARY OF THE MODEL

The model implements the DOCSIS architecture defined in [1]. Packets sent over the downstream channel are broken into 188 byte MPEG frames each with 4 bytes of header and trailer. The model accounts for physical layer overhead including framing bits and forward error correction data. The downstream channel supports an optional token bucket-based service rate. Each SID service queue is treated in a first come first serve manner. Depending on traffic dynamics, queueing can occur at either the SID queue or the downstream transmission queue. The maximum size of either queue is a simulation parameter.

All CMs receive periodic MAP messages from the CMTS that identify future upstream scheduling opportunities over the next MAP time. If provisioned with a periodic grant, a CM can send at its next data grant opportunity. For best effort traffic, a CM must request upstream bandwidth from the CMTS using a contention-based mechanism. To improve efficiency, a CM can request bandwidth to transport multiple IP packets in a single DOCSIS frame by issuing a concatenated request. Further, a CM can piggyback a request for bandwidth on an upstream data frame. If a CM receives a grant for a smaller number of minislots than were requested, the CM must fragment the data to fit into the assigned slots. Our model supports concatenation, piggybacking and fragmentation.

Figure 1 illustrates the MAP layout used in our model. The first slot at the left of the MAP represents time 0 in the MAP time. Data slots are placed at the beginning of the MAP and contention slots are placed at the end. Figure 2 illustrates the upstream transmission of a 1500 byte IP datagram from a TCP source directly connected to a CM to a sink connected to the CMTS. In Figure 2, time progresses in the downwards direction. We assume collisions do not occur. Assuming a MAP size of 80 slots, an upstream channel capacity is 5.12Mbps and there are 4 ticks per slot, 96 slots are required to transport the entire packet. The small dark square box positioned at the beginning each MAP time in the figure represents the transmission of the MAP message in the downstream direction. Our model sends the MAP at the beginning of each MAP time. Each MAP describes the slot assignments for the next MAP time. The IP packet arrives at the CM during the j 'th MAP time at time T-0. The CM sends the bandwidth request message at time T-1 and receives the data grant at time T-2. The grant is allocated in the $j+2$ MAP time. The CM sends the frame at Time T-3 and is received by the

CMTS at time T-4. The time between T-3 and T-0 is the access delay which represents the total time a packet is delayed over the DOCSIS network not including transmission or propagation time. The model can be configured to allocate a specific number of contention request slots each MAP. Or, in addition to a minimum number of contention re-

quest slots, all unused slots can be designated for contention requests.

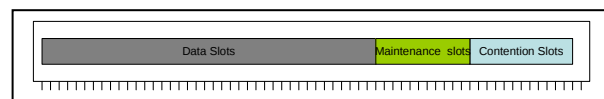


Figure 1. MAP layout

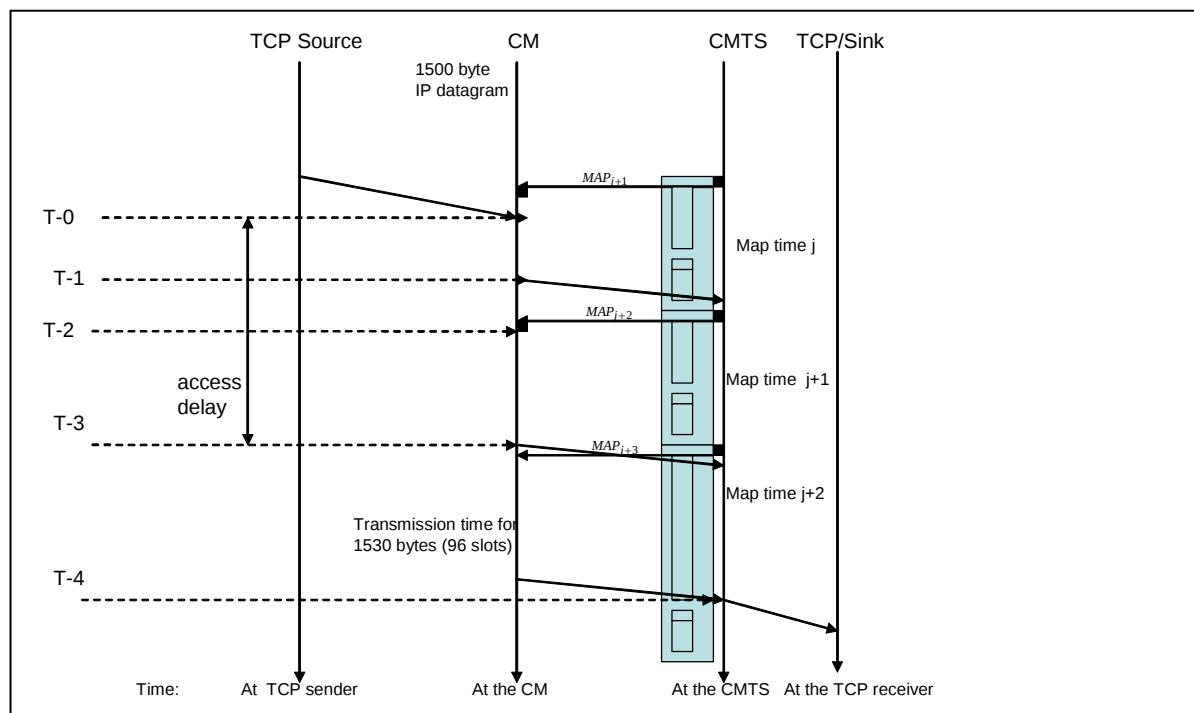


Figure 2. Upstream operation

IMPACT OF DOCSIS ON TCP

The results we report were based on simulation experiments using the network shown in Figure 3. The DOCSIS parameters were based on optimal configuration parameters that we found in a previous study [3]. A set of user nodes were attached to the CMs and a set of server nodes were located in the wired network. The traffic generators utilized realistic traffic models consisting of a combination of web, P2P and streaming traffic. The network and web traffic models were based on the “flexbell” model defined in [4]. In addition to downstream web traffic, we configure 5% of the CMs to generate downstream low speed UDP streaming traffic (i.e., a 56Kbps

audio stream), 2% of the CMs to generate downstream high speed UDP streaming traffic (i.e., a 300Kbps video stream) and 5% of the CMs to generate downstream P2P traffic. The P2P model (based on [5]) incorporates an exponential on/off TCP traffic generator that periodically downloads on average 4Mbytes of data with an average idle time of 5 seconds between each download. The downstream transmission queue at the CMTS was configured to hold a maximum of 50 packets. We limited the number of packets that can be concatenated in a single frame to two. The DOCSIS and Web traffic simulation parameters are shown in Figure 4.

We varied two parameters in the experiments, the MAP_TIME and the number of CMs. For a given MAP_TIME setting, we varied the number of CMs from 100 to 500. We do this for six MAP_TIME settings ranging from .001 to .01 seconds.

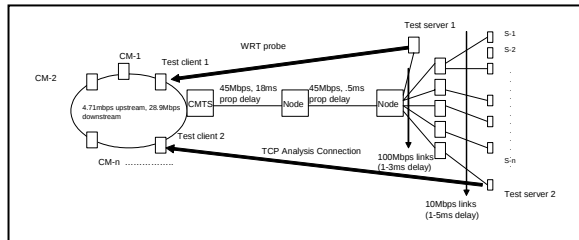


Figure 3. Simulated network

We obtained the following statistics for each run:

Collision rate: Each time a CM detects a collision it increments a counter. The collision rate is the ratio of the number of collisions to the total number of upstream packets transmissions attempted.

Downstream and upstream channel utilization: At the end of a run, the CMTS computes the ratio of the total bandwidth consumed to the configured raw channel bandwidth. The utilization value reflects the MAC and physical layer overhead including FEC bits.

Average upstream access delay: All CMs keep track of the delay from when an IP packet arrives at the CM in the upstream direction until when it actually gets transmitted. This statistic is the mean of all of the samples.

Web response time: a simple TCP client server application runs between test client 1 and the test server 1. Test server 1 periodically sends 20Kbytes of data to test client 1. With each iteration, the client obtains a response time sample. The iteration delay is set at 2 seconds. At the end of the test, the mean of the response times is computed. The mean web response time (WRT) can be correlated to end user perceived quality by using a very coarse rule of thumb that says end users are

bothered by lengthy download times when the mean WRT metric value exceeds 1 second. We do not claim this to be an accurate measure of end user quality of experience. Instead, it is a convenient, reproducible performance reference.

Model Parameters

Upstream bandwidth 5.12Mbps
Preamble 80 bits
Downstream bandwidth 30.34Mbps
4 ticks per minislot
Default map time: 2 milliseconds (80 minislots per map)
Fragmentation Off, MAP_LOOKAHEAD = 255 slots
Concatenation ON
Backoff Start: 8 slots, Backoff stop: 128 slots
12 contention slots, 3 management slots
Simulation time: 1000 seconds

Web Traffic Model Parameters

Inter-page: pareto model, mean 10 and shape 2
Objects/page: pareto model, mean 3 and shape 1.5
Inter-object: pareto model, mean .5 and shape 1.5
Object size: pareto model, mean 12 (segments) shape 1.2

Figure 4. Simulation parameters

Web Congestion Experiment Results

Figures 5a and 5b plot the channel utilization as the load increases. The downstream utilization reaches a maximum of about 64% with a MAP_TIME setting of .001 second. In this case, 12 contention slots per MAP is sufficient. For larger MAP_TIME values, the downstream utilization ramps up to its maximum value and then decreases at varying rates as the load increases. As the collision rate grows, downstream TCP connection throughput decreases. Limiting each MAP to 12 contention slots results in fewer total contention request opportunities as the MAP_TIME grows. This explains the high collision rates and reduced downstream utilization for the runs with large MAP_TIME settings.

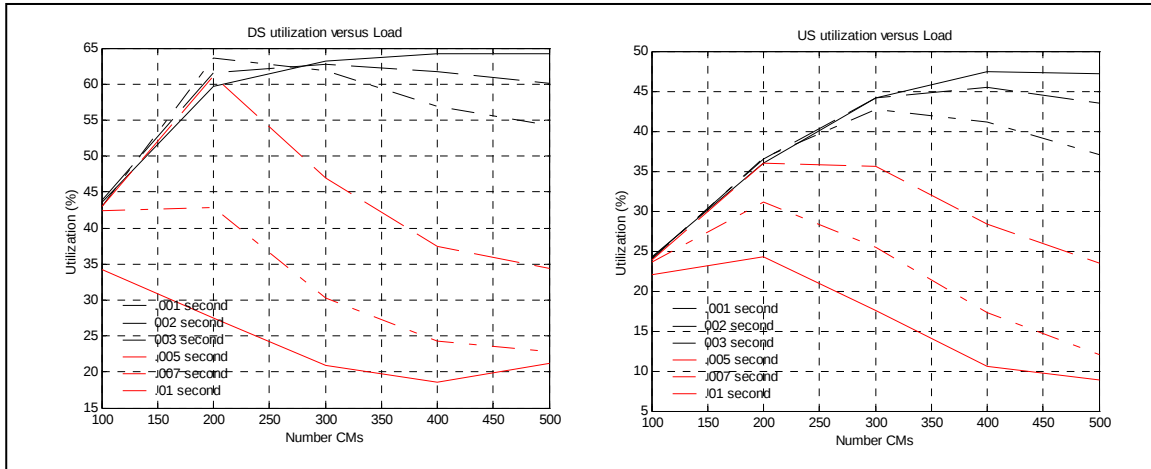


Figure 5a. Downstream channel utilizations

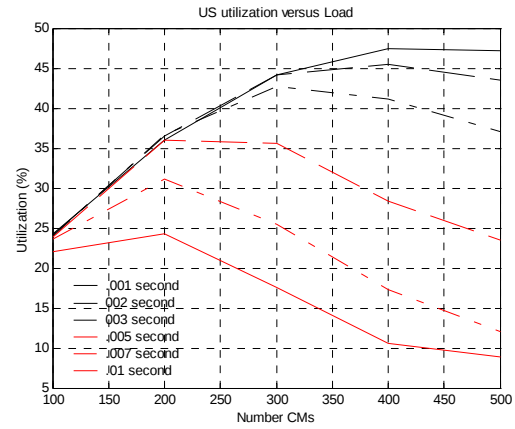


Figure 5b. Upstream channel utilizations

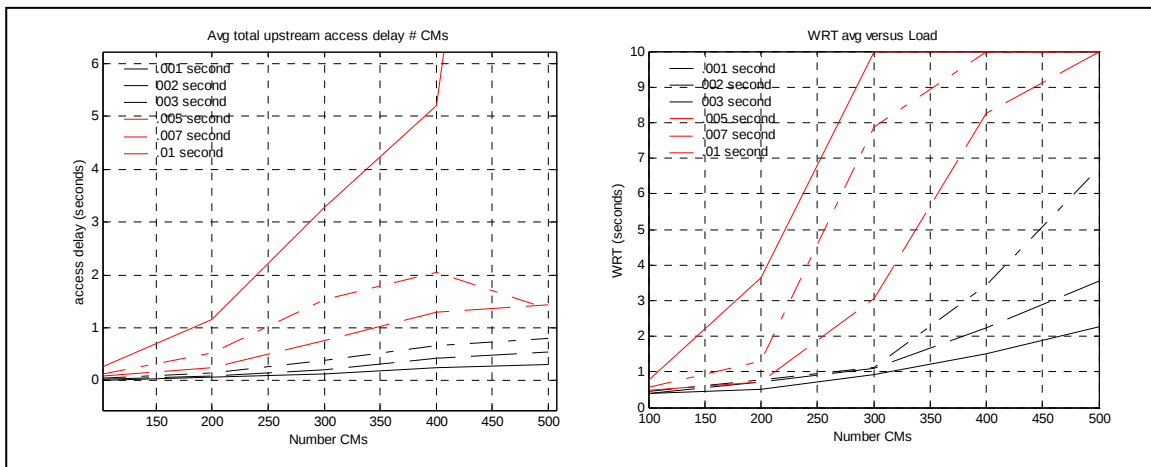


Figure 6a. Upstream access delay

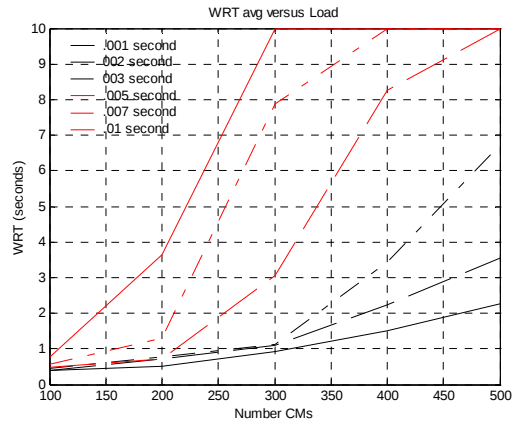


Figure 6b. Web response time metric results

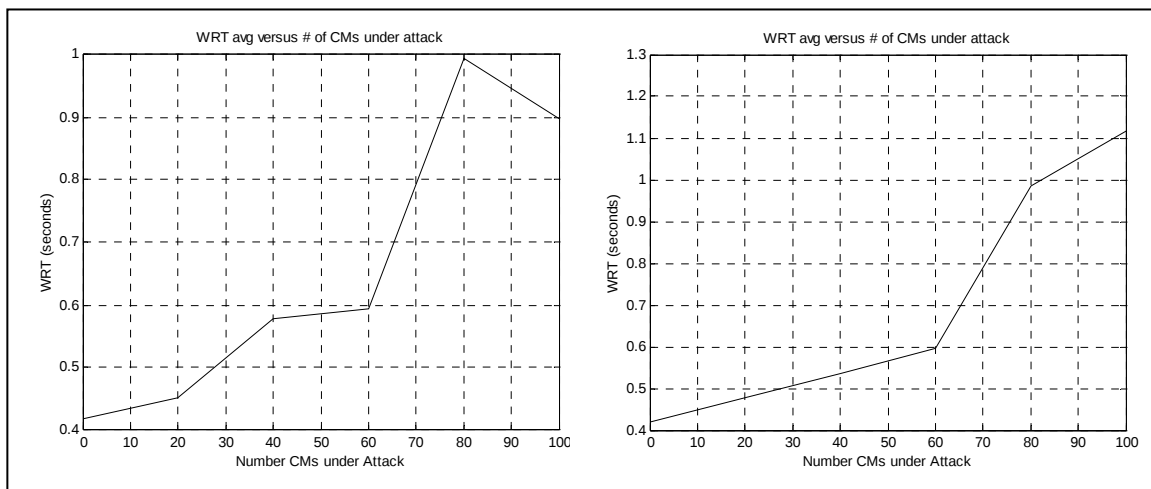


Figure 7a. WRT results without rate control

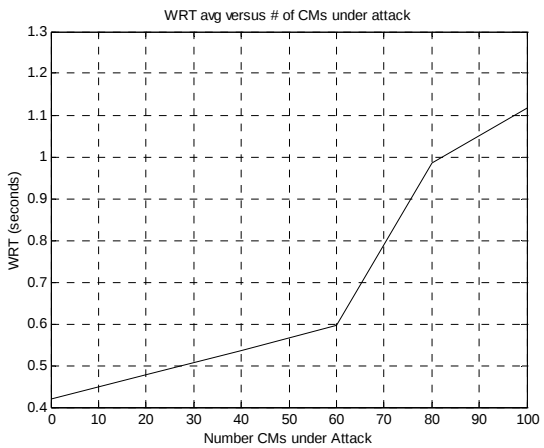


Figure 7b. WRT with 2Mbps DS rate control

Figure 6a shows that the average upstream access delay becomes very large at high loads when configured with large MAP_TIME settings. Even for lower MAP_TIME values, the access delay was significant. For a MAP_TIME of .002 seconds, the access delay exceeded .5 seconds at the highest load level. To assess the impact of the cable network on end-to-end performance we monitored web response times. Using the rule of thumb described earlier, Figure 6b suggests that for MAP_TIME settings less than .005, up to 300 users can be active before performance becomes bothersome to end users.

When 100 users are active, the collision rate is about 50%. What makes this result alarming is that the web traffic model accounts for the heavy tailed distribution associated with web user idle times. Consequently, the number of users actually competing for bandwidth at any given time is much less than 100. As the load increased, the collision rate approached 90% depending on the MAP_TIME setting.

When the dominant application is web browsing, the majority of data travels in the downstream direction. At high loads, the network can become packet rate bound causing ACK packets accumulate in the CM upstream queues waiting for transmission opportunities. Piggybacking is of limited benefit since ACKs that arrive back-to-back are sent in a concatenated frame. Concatenation can be helpful although it drastically increases the level of ‘ACK compression’ experienced by downstream TCP data transfers [3]. ACK compression occurs when a network causes TCP acknowledgement packets to ‘bunch’ at some point leading to bursty TCP send behavior which in turn contributes to higher loss rates and poor network utilization.

We repeated the study using different parameters and features of the model. The results are virtually identical if we turn on a downstream

service rate of 2Mbps, if we turn on ACK filtering or if we allocate all unused slots for contention requests. However, if we increase the downstream transmission queue size at the CMTS (the point where loss occurs) from 50 to 300 packets, loss no longer occurs and downstream utilization approaches 75%. While a larger buffer improves performance, the important result is that the downstream TCP traffic is subject to extreme levels of ACK compression caused by DOCSIS.

DoS VULNERABILITIES

In this section we show that it is possible for a hacker to take advantage of the inefficient contention-based upstream bandwidth allocation process by initiating a denial-of-service (DoS) attack. To accomplish the DoS attack, a host located outside the DOCSIS network must learn the IP address of a number of CMs that share the same upstream channel. The attacker transmits either a ping or a TCP SYN packet to the targeted CMs at a frequency that depends on how many CMs are under attack. The objective of the attack is to cause a large number of contention-based requests resulting in high collision rates and subsequently poor network performance. This type of attack has been identified in 802.11 networks where an attacker stimulates stations to initiate RTS/CTS exchanges leading to dramatically reduce network efficiency [15].

We simulate an attack using the network model illustrated in Figure 3. The configuration was identical to that described in Figure 4. We set the MAP_TIME to .002 second. There were 100 CMs but the number of CMs under attack was varied. The collision rate increased from 48% to 68% as the number of CMs under attack increased from 0 to 100. The downstream utilization dropped from 45% to 10%. Figure 7a shows that the web response times increased by a factor of 3. The web response time monitor was located at a CM that was not under attack (i.e., test client

1 in Figure 3). In a separate experiment, we included the test client 1 CM in the attack and found that the CM was not able to complete a single web response time sample. Not surprisingly, a CM subject to a flooding attack effectively makes the access network unavailable to the subscriber.

We ran the denial-of-service experiment a second time with downstream service rates set to 2Mbps. The results were virtually identical to the previous results. Figure 7b shows the average web response times from test client 1 when this node was not under attack also increased by almost a factor of 3. The result suggests that a 2Mbps downstream service rate will not protect the network from the attack.

FAIRNESS ISSUES

There are several fairness issues that can arise in a DOCSIS network primarily caused by upstream packet rate limitations. The first issue is that DOCSIS exhibits bias against TCP connections running over paths with small MTU sizes. If two CMs are each transporting data from separate but identically configured TCP connections with the exception that one connection has a negotiated MSS of 512 bytes and the other connection uses an MSS of 1492 bytes, the connection that generates the larger packets will consume more bandwidth than the other connection.

A second issue, applicable to the downstream direction, involves streaming video protocols, such as TFRC, that claim to be TCP-friendly. Because DOCSIS systems can be packet rate bound, rate-based protocols such as TFRC that do not require an ACK stream to clock new data can consume larger amounts of bandwidth than comparable TCP connections.

To demonstrate this second issue, we modified the previous web scenario experiment by adding an FTP-like TCP flow between one of

the CMs and a server and a similar FTP-like TFRC flow between another CM and server pair. The MAP_TIME was .002 seconds and the number of contention request slots per map was set to 12. We performed six runs, increasing the number of CMs from 0 to 500. Figure 8 plots the TCP and TFRC connection throughputs for each run. The TFRC flow obtains roughly 3-7 times the bandwidth of the TCP flow depending on the number of CMs. When there are just the two CMs competing (this is the 0 point on the x-axis of Figure 8), the TFRC and TCP flows achieve a throughput of 18Mbps and 6 Mbps respectively. The TFRC flow by itself (i.e., if we do not run the competing TCP connection) obtains about 22Mbps while the TCP flow by itself obtains about 12.5 Mbps. If the channel bandwidths increase, the maximum TCP throughput does not change (because TCP throughput is packet rate limited in the upstream direction rather than limited by downstream bandwidth). The TFRC flow does not have this limitation and can consume higher downstream bandwidths.

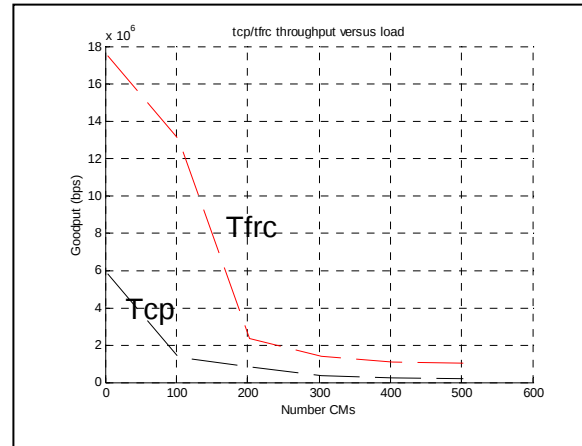


Figure 8. TCP and TFRC throughput

BANDWIDTH MANAGEMENT

In our current research, we are exploring bandwidth management to address these issues. One component of our work is to develop a protocol aware scheduling algorithm that predicts future CM upstream bandwidth needs and provides unsolicited grants. A sec-

ond component is to develop bandwidth management algorithms that manage bandwidth based on a particular policy or service. For example, as an alternative to a pay-per-use policy, a provider might desire a policy where subscribers that consume large amounts of bandwidth in either the upstream or downstream directions are ‘punished’ by being placed in a state of reduced service rates for a given time period.

We have prototyped such an algorithm in our simulation model. The motivation for the algorithm is that future cable services will offer much higher service rates, possibly on the order of tens of megabits per second. To manage fairness issues or to facilitate new service options, dynamic bandwidth management is required. The objective of our algorithm is to prevent the large number of well behaved subscribers from adverse affects caused by a few high bandwidth users (referred to as ‘heavy-hitters’). The algorithm has three components: detecting poor quality of service observed by normal user, identifying the heavy-hitters and regulating the heavy users to solve the problem.

The algorithm runs at the CMTS and does not require any changes at the CM nodes. The algorithm can be used on the downstream channel or the upstream channel (or both). We have applied the algorithm to manage downstream bandwidth. The majority of subscribers are well behaved in the sense that they consume a reasonable amount of bandwidth over large time periods. A few subscribers are not well behaved (i.e., the heavy-hitters) and consume a disproportionate amount of bandwidth over large time scales. To simplify the discussion, we assume that there is one user per CM.

Detecting poor quality of service observed by normal user

We use the WRT metric described earlier to characterize the quality of service perceived

by a subscriber. A node attached to a CM node periodically sends a request to an HTTP server for a 20Kbyte object. The time taken for this download is monitored periodically and is averaged over a time-scale defined by the parameter *WRTM* (WRT monitor interval). We assume that when the average of the WRT samples over a *WRTM* time period approaches 1 second users will perceive poor quality. Once this situation is detected, the algorithm identifies the heavy-hitters that are contributing to congestion in the network.

Identifying heavy-hitters in the cable network

The algorithm maintains the average bandwidth rate (*ABR*) of all active users. The time interval over which the rate is averaged is defined by the parameter *TAVG*. An active user is a user whose *ABR* is not zero over a *TAVG* amount of time. Based on the maximum channel capacity and the number of active users present, a fair bandwidth rate (*FBR*) of each user is calculated using the following equation:

$$FBR = (Maximum\ channel\ capacity) / (Number\ of\ Active\ users).$$

The number of active users present in the network is based on samples averaged over a period of time defined by the parameter *TNUS* (time for number users sample). Any user whose *ABR* is above a threshold based on the *FBR* is considered to be a heavy-hitter. The threshold value is represented by the parameter *THUSR*.

The timescale parameters, *TAVG* and *TNUS*, allows a cable service provider to implement different policies. For example, a small *TAVG* on the order of minutes, can be used to ensure that TFRC flows consume a fair share of bandwidth. A cable service provider might want to detect users who operate servers (e.g., peer-to-peer or web servers). This can be handled by setting the *TAVG* to days. What makes the algorithm unique is the fact that a

heavy-hitter is not punished unless it is impacting other users. The extent of the punishment is determined by algorithm parameters.

Regulating the heavy users

Once a heavy-hitter has been identified, the next step is to regulate it to improve network performance. We implemented the policy that heavy-hitters never get more than the fair share of the bandwidth. The rate regulation continues for a configurable period of time (*TREG*). The rationale for ‘punishing’ the heavy-hitters by limiting their bandwidth to the fair share is to ensure that they no longer impact well-behaved users. Since we only consider the number of active users in calculating the fair share, it is possible that the channel might be under-utilized. For instance, assume there are 100 users using 30Kbps bandwidth. The fair share will be 300Kbps for a 30 Mbps channel. While being punished, a heavy-hitter can consume a maximum of 300Kbps even though additional bandwidth might be available.

Simulation verification

We demonstrate the algorithm using the simulation network in Figure 3. We configured 154 CMs to generate the traffic mix described in Figure 4. We configured 6 additional CMs to maliciously consume large amounts of downstream bandwidth using UDP traffic sources. All the heavy-hitters were started and stopped at the same time. When the simulation starts, the 100 web users are started. The heavy-hitters start at around 3000 seconds collectively generating around 35Mbps of traffic. Figure 9 plots the aggregate downstream bandwidth over the experiment. At time 3000 seconds we see the aggregate bandwidth increase as the heavy-hitters start. The algorithm smoothly adapts subscriber rates to the penalized value. The subscriber will be in the penalty state for about 4 hours. All web traffic stops time 13500 except for traffic generated by the 6 heavy-hitters. As the

algorithm detects available bandwidth, it allocates more bandwidth to the heavy-hitters. If there are no other users, the fair share allocated to the heavy-hitters will consume all available bandwidth. More likely there will be other users in which case the *FBR* will limit the heavy-hitters but not affect well behaved users.

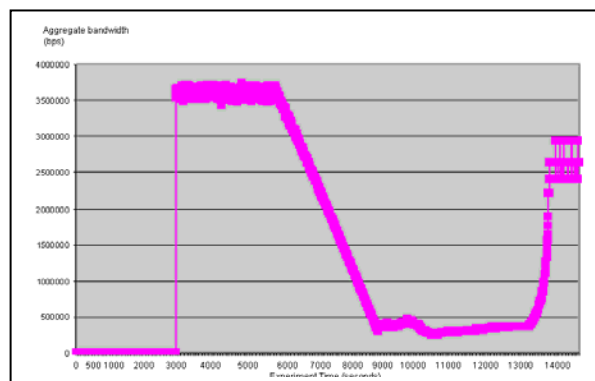


Figure 9. Bandwidth management algorithm

RELATED WORK

While the intent of the IEEE’s 802.14 effort was to provide ATM services over a hybrid fiber coaxial (HFC) medium, the operation of the MAC layer is similar to that supported by DOCSIS. Therefore, prior 802.14 research is relevant. The work in [6] found that TCP throughput over an 802.14 network is low primarily due to ACK compression. The authors propose two solutions: one involving piggybacking and a second involving TCP rate smoothing by controlling the ACK spacing. The authors found that piggybacking can help reduce the burstiness associated with the ACK stream in certain situations. However it is limited in its abilities to effectively match offered load over a range of operating conditions. The author’s second solution is to control the TCP sending rate by measuring the available bandwidth and calculating an appropriate ACK rate and allowing the CM to request a periodic grant that provides sufficient upstream bandwidth to meet the required ACK rate. We distinguish our work by focus-

ing on the latest DOCSIS standards (1.1 and 2.0) and using more realistic traffic loads.

The observation in [7] is that an HFC network presents difficulties for TCP due to the asymmetry and due to high loss rates (possibly as high as 10-50%). Due to the problems of TCP/Reno in these environments[8,9,10], the authors propose a faster than fast retransmit operation where a TCP sender assumes that a packet is dropped when the first duplicate ACK is received (rather than the usual triple duplicate ACK indication). The motivations behind [7] are not relevant with the latest DOCSIS standards as DOCSIS 2.0 provides nearly symmetric access links with low packet loss rates as long as the plant is well engineered.

The performance of TCP over asymmetric paths has been thoroughly studied [11,12,13]. A network exhibits asymmetry with respect to TCP performance if achieved throughput is not solely a function of the link and traffic characteristics of the forward direction but in fact depends on the impact of the reverse direction. Most of the prior work was focused on highly asymmetric paths with respect to bandwidth where the normalized asymmetry level (i.e., the ratio of raw bandwidths to the ratio of packet sizes in both directions) typically would be on the order of 2-4 [11]. In DOCSIS, depending on the service rate configuration, the level of bandwidth asymmetry is small (or nonexistent). Instead, DOCSIS exhibits packet rate asymmetry due to low upstream packet rates with respect to downstream capacity. However the problem symptoms are similar. Various methods have been proposed to alleviate the TCP over asymmetric path problems including header compression and modified upstream queue policies(drop-from-front, ACK prioritization, ACK filtering) [11,12,13,14]. Some of these ideas can be applied to DOCSIS. For example, a CM that supports ACK filtering could drop 'redundant' ACKs that are queued. We

have implemented this and found that while it does increase the acknowledgement rate, it also increases the level of ACK compression. ACK reconstruction could be implemented in the CMTS to prevent the increased level of ACK compression from affecting performance. We plan on addressing this in the future.

CONCLUSIONS

Using simulation we have identified several issues. First we saw that DOCSIS can affect the ACK stream in the upstream direction resulting in bursty downstream dynamics. Second, we have identified a possible DoS vulnerability in DOCSIS. Taking advantage of the inefficiency associated with upstream packet transmissions, a hacker can negatively impact network performance by periodically stimulating (e.g., by ping or TCP SYN packets) a number of CMs at a frequency that depends on the number of CMs under attack. The signature for this attack would be different than that of traditional flooding attacks as the amount of bandwidth consumed in the downstream direction is low. Finally, we illustrated that a *TCP-friendly* protocol turns out to be *TCP-unfriendly* in a DOCSIS environment because the model of TCP behavior incorporated by TFRC fails to accurately capture how TCP performs in a DOCSIS environment.

We presented an algorithm that is designed to help manage user traffic in DOCSIS networks. While the algorithm might not be appropriate for today's networks that rely on low service rates or that involves penalties for bandwidth misuse, our work is intended for future higher speed cable access networks that are likely to offer service rates on the order of tens of Mbps. In these environments intelligent bandwidth management will be required.

REFERENCES

1. Cable Television Labs Inc. , CableLabs, "Data-Over Cable Service Interface Specifications- Radio

Frequency Interface Specification”, SP-RFIV2.0, available at <http://www.cablemodem.com/specifications/specifications20.html>.

2. The Network Simulator. Available at : <http://www-mash.cs.Berkeley.EDU/ns/>.

3. J. Martin, N. Shrivastav, “Modeling the DOCSIS 1.1/2.0 MAC Protocol”, Proceedings of the 2003 International Conference on Computer Communications and Networks”, Dallas TX, October 2003.

4. A. Feldmann, et. Al., “Dynamics of IP Traffic: A study of the role of variability and the impact of control”, SIGCOM99.

5. S. Saroiu, P. Gummadi, S. Gribble, “A Measurement Study of Peer-to-Peer File Sharing Systems”, Multimedia Computing and Networking (MMCN), Jan 2002.

6. R. Cohen, S. Ramanathan, “TCP for High Performance in Hybrid Fiber Coaxial Broad-band Access Networks”, IEEE/ACM Transactions on Networking, Vol. 6, No. 1, February 1998.

7. O. Elloumi, et. Al., “A Simulation-based Study of TCP Dynamics over HFC Networks”, Computer Networks, Vol. 32, No. 3, pp 301-317, 2000.

8. O. Elloumi, et. Al., “Improving Congestion Avoidance Algorithms in Asymmetric Networks”, IEEE ICC 97, June 1997.

9. K. Fall, S. Floyd, “Simulation-based Comparisons of Tahoe, Reno and SACK TCP”, CCR, Vol 26, No. 3, July 1996.

10. J. Hoe, “Improving the Startup Behavior of a Congestion Control Scheme for TCP”, SIGCOMM 96, August 1996.

11. H. Balakrishnan, et. Al., “The Effects of Asymmetry on TCP Performance”, ACM/IEEE International Conference on Mobile Computing and Networking, Sept. 1997.

12. T. Lakshman, U. Madhow, B. Suter, “Window-based error recovery and flow control with a slow acknowledgement channel: a study of TCP/IP performance”, INFOCOM97, April 1997.

13. V Jacobson, “Compressing TCP/IP Headers for Low-Speed Serial Links”, Feb 1990, RFC 1144.

14. L. kalampoukas, A Varma, K. Ramakrishnan, “Improving TCP Throughput over Two-Way Asymmetric Links: Analysis and Solutions”, SIGMETRICS 98, June 1998.

15. Saikat Ray, Jeffrey B. Carruthers, and David Starobinski, “RTS/CTS-induced congestion in ad-hoc wireless LANs,” in IEEE Wireless Communication and Networking Conference (WCNC), March 2003.

DOWNLOADABLE SECURITY

James William Fahrny
Comcast Cable Communications

Abstract

This paper will define a common security architecture that overcomes some of these obstacles and issues described in the background section above. This common hardware security platform can be used to secure Broadcast Conditional Access systems, Video On Demand services, Digital Rights Management, and the Authorized Service Domain services in extending CA into the home network.

This proposed paper would cover the following topics:

- *Architecture Block Diagram of the Downloadable Security System*
- *Description of the secure download mechanism and how it can be secured.*
- *Analysis of how this advances security for video and audio content*
- *Definition on how this can be used to perform "Hardware Renewability" with a software download using FPGA technology.*
- *Describes how the paradigm of revocation and renewability are modified for a better customer experience.*

Analysis of how the downloadable features can be applied to various applications of the hardware platform including CAS, VOD, DRM, Trusted Domain, Streaming and Personal Computing

BACKGROUND

Conventional implementations of media (e.g., video, audio, video plus audio, and the like) program stream delivery systems (e.g., cable, satellite, etc.) include a head-end where the media programming originates (i.e., is encoded and compressed, groomed, statmuxed, and otherwise appropriately processed), a network (e.g., cable or satellite) for delivery of the media programming to the client (i.e., customer, user, buyer, etc.) location, at least one set top box (STB) at the client location for conversion (e.g., decryption and decompression) of the media programming stream, and at least one respective viewing device such as a television (TV) or monitor that is connected to the STB. Alternatively, the STB may be eliminated, and decryption and decompression may be implemented in the receiving device.

Conventional head-ends and STBs employ particular matching encryption/decryption and compression/decompression technologies. However, there is little standardization of particular matching encryption/decryption across media program stream delivery system vendors. The encryption/decryption and compression/decompression technologies in the particular conventional system are fixed and often proprietary to the vendor. Furthermore, conventional media service processing and delivery systems typically implement security processes in connection with individual implementations of point of deployment, CableCard, Smartcard, etc. systems.

Transitions to upgrades in encryption/decryption and compression/decompression technologies are, therefore, expensive and difficult for the media program stream delivery system vendors to implement. As such, customers can be left with substandard service due to the lack of standardization and the reduced competition that the lack of standardization has on innovation in media service delivery. The lack of standardization also restricts the ability of media service providers to compete. For example, customers may have viewing devices that could take advantage of the improved technologies; however, media stream delivery system upgrades may be impossible, impracticable, or not economically feasible for vendors using conventional approaches. A significant level of customer dissatisfaction or vendor cost may result and the ability of media service providers to improve service and/or add new services is greatly restricted.

Thus, it would be desirable to have a system and a method for CA download and reconfiguration that overcomes the deficiencies of conventional approaches.

SUMMARY

This paper generally describes an improved system and method for security processing digital media streams. The improved system and method for security processing media streams of the present invention may be compatible with previously used (i.e., legacy) systems and methods using all levels of media stream processing and delivery service (i.e., basic to high-end) as well as adaptable to future implementations, and that is flexible, renewable, re-configurable, and could support simultaneous multiple security systems and processes.

This paper proposes a method of multi-stream security processing and distributing digital media streams. The technology comprises generating encrypted digital media streams. The method further comprises coupling a network to the head-end and receiving the encrypted digital media streams at the network. The technology yet further comprises coupling at least one receiver to the network and receiving the encrypted digital media streams at the receiver, and presenting a decrypted version of the encrypted digital media streams using the receiver. At least one of the head-end and the at least one receiver comprises a security processor that may be configured to provide at least one of simultaneous multiple encryption and simultaneous multiple decryption processing of the digital media streams.

This paper describes a system for multi-stream security processing, key management, and distributing digital media streams, a security processor configured to provide at least one of simultaneous multiple media transport stream decryption and encryption processing is provided. The single chip solution described in this paper includes a security processor, a controller and a plurality of digital decryption engines. The digital decryption engines may be selectively parallel coupled by the controller for simultaneous operation in response to a predetermined security configuration.

Though this paper describes a future vision Security System On a Chip (SSOC), the current technology widely deployed is done with separate physical devices. The Large Scale Integration (back-end) device contains all codecs, transport functions, de-compression, general purpose processor, memory, and other subsystems. The

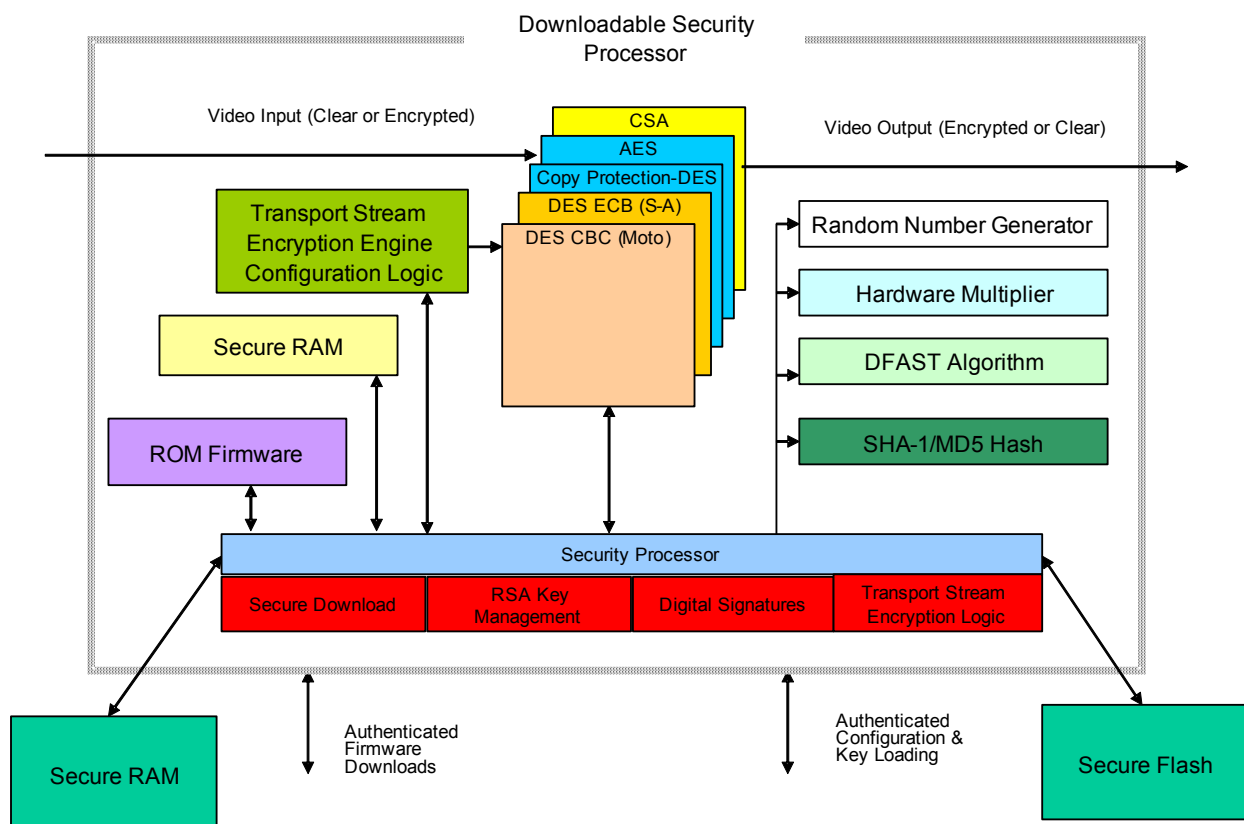
Security Processor is a separate chip since it typically has on-board flash memory and numerous layers of tamper resistance and countermeasures to prevent hacker attacks. Based on current device fabrication technologies, the SSOC is not the most cost effective solution though it can be made more secure. In the future, there may be technology that enables the SSOC in a cost effective manner by including or replacing all of the tamper systems and countermeasures on the larger device.

ARCHITECTURE

The following diagram defines the elements of a downloadable and re-configurable security processing system. This diagram includes the key management

system and secure down load system to install a new key management system. In addition, the re-configurable security transport system is defined as part of the overall technology.

The following diagram shows a single block diagram, which is a logical representation. The transport stream decryptors can be packaged in a separate integrated circuit (IC) with the other set-tops functions like decoders, demux, graphics engines, etc. When the system is packaged separately, there must be methods implemented to secure communications between the transport stream decryptors and the security processor in the client device. This security is not in the scope of this paper.



Authentication

When devices are installed on a cable network for the first time, a discovery process must occur. The Head-end server would broadcast an announcement message much like a DHCP server. The client will respond to the “announce” message with credentials to authenticate the client security processor. The head-end server will then present its credentials to the client security processor so that the client device can trust the server and the server can trust the client security processor. Credentials are presented with digital signatures for authenticity and the public key is used to verify the credentials on the receiving device.

Download: Obtaining a client’s network personality

Once the server trusts the client security processor and the client security processor trusts the server, the security processor determines whether it needs a personality in the form of a security client depending on the network information that it receives in the authentication process. If a download is required, a secure key exchange occurs to setup for the transfer of an encrypted and digitally signed security image to the security processor in the client device. Once the client is validated through signature verification and decryption, it is loaded into the security processor and executed.

This process permits a security processor to obtain its security personality (Scientific-Atlanta, Motorola, NDS, Nagra, etc) when it checks into securely the first time. This technology also allows a device to move from one network to another whereby the client will determine that it has incorrect security client software loaded, and deletes it from memory. At this time, the client security

processor will request the head-end server to download a security client for the new network personality. Finally, the secure download process can be used to simply upgrade the key management methods while preserving entitlements, purchases, credits and other important data stored in the security processor.

SECURING THE DOWNLOAD

The anchor of security and trust within this technology is in the security of the download. The client should ONLY be able to download new firmware when the head-end server commands the client to receive a new download. The client should not be able to force a download outside of the proper network personality changes. If any of this is incorrectly design, the overall security is subject to severe compromise. To accomplish this strong security, the head-end must “unlock” the download ability of the client security processor. If the client is locked, the security processor cannot be loaded with a new client.

In the same way, the image being sent to the client security processor must not be tampered and likely has elements of data that should not be viewed. This leads to the use of digital signatures and symmetric encryption of the image. The signature protects the image from being modified and the encryption protects data elements from being viewed.

ADVANCING CONTENT SECURITY

In the current systems, there are typically no methods to upgrade the system for fear that tampering or countermeasures would be more easily installed. Therefore, most components in the security architectures are constructed so that they cannot be modified.

This can be great from a security view but leaves no ability to adapt to the changing world of content and content delivery systems.

The system proposed in this paper is not as static with the secure download mechanism and therefore creates new abilities in support of potential business models as they are developed. Similarly, the longevity of a renewable security system that can adapt over time but remain secure appears to be greater than conventional security methods.

USING SECURE DOWNLOAD FOR HARDWARE RENEWABILITY

There is another unique development that recently presented itself in the development of this technology. Field Programmable Gate Arrays (FPGA) have been used LSI hardware components where one can develop hardware logic and load the logic language into an FPGA to achieve dynamic hardware. This technology has been very expensive in the past.

However, recently IBM and other research facilities have developed FPGA technology in 90 nanometer geometries of chips that is extremely cost effective. One could effectively include 4,000 to 10,000 gates of FPGA into a security processor and leave it blank for future use. If a new algorithm is required because something is compromised, the FPGA could be used to create a hardware accelerator of a new algorithm. In this case, all of the client security processors would be downloaded with hardware logic that would be loaded into the FPGA section of the security processor to enable the new algorithm.

CHANGING THE PARADIGM: RENEWABILITY INSTEAD OF REVOCATION

One of the largest problems in security systems is that of revocation. Revocation is very operationally unfriendly to manage and is really built in a manner to not have a large scale system revocation of keys. The single biggest issue with revocation is that a revocation event typically disables legitimate customers experience.

In this Downloadable Security system, we propose that the renewability be used in place of revocation in all cases possible. Since this system can securely transport data from the head-end server to the client security processor, keys of many types (authentication and encryption) can be renewed when compromised or periodically if desired. Clearly, if keys are renewed in a “live” system, synchronization of the transition must be managed using a solid time base so that the customer experience is completely uninterrupted.

In any case, this system provides a much better possibility with renewal since the compromised systems are not actively shutdown but is passively allowed to expire when the crypto period of their entitlements end. Paying customers are then renewed in this process to a new key set so that their experience is uninterrupted.

FUTURE APPLICATIONS

This technology was developed to focus on a certain set of problems in the Broadcast Conditional Access domain. However, after further review, this technology will be very effective if applied to On Demand security, Home Network Content distribution, Trusted Computing Platforms, Digital Rights Management, and Interactive Gaming.

In fact, this technology is so flexible that one could deploy a product with the security processor hardware and a specific security application or profile. Later the security processor could be upgraded to add security management for one of the other technologies as it is added as a service to the network.

For example, a system could be deployed with broad CAS and later be upgrade to support VOD, or Home Network Content security with a secure upgrade in the field.

CONCLUSIONS

To summarize, we believe that the Secure Download technology described in this paper

will provide the next generation of security for Broadcast, Video On Demand, and Streaming media systems. If the Large Scale Integrated (LSI) devices are design with some flexibility for the future, this system will have a tremendous longevity and a strong ability to counter any hacker attacks to steal services or clone devices in the field for signal theft.

In the same way, we believe that the usage of this technology will grow with time since we are only viewing the initial stage of this new paradigm at the present time. Applying this technology to Home Networking, IP video delivery systems or even Digital Rights Management will increase greatly over the next couple of years.

HIGH SPEED MULTIMEDIA HOME NETWORKING OVER POWERLINE

Haniph A. Latchman¹, K. Afkhamie³, S. Katar³,
R. E. Newman², B. Mashburn³, L. Yonge³

¹ECE Department, University of Florida, Gainesville FL 32611

²CISE Department, University of Florida, Gainesville, FL 32611

³Intellon Corporation, 5100 W. Silver Spring Blvd., Ocala, FL, 34482

Abstract

This white paper describes the unique challenges associated with high speed digital communication over existing in-building powerlines. The solutions provided by the 14 Mbps HomePlug 1.0 protocol are described and an overview of the 200 Mbps HomePlug AV protocol is given. The latter protocol is optimized for multimedia voice and video services, while also providing high speed data communication.

INTRODUCTION

Interest in Powerline Communications

There has been a great deal of recent interest in leveraging the existing electrical wiring within and connected to buildings for high speed digital communications [1]. In-home LANs using powerline communication (PLC) are now a reality with products based on the HomePlug 1.0 standard in use worldwide since 2000. [2][3]. PLC LANs using the 14 Mbps HomePlug 1.0 chipsets, provide full house coverage at typical TCP data rates of 5-7 Mbps, and exhibit greater stability than competing wireless LAN solutions [4,10].

In addition there is current activity in the deployment of Broadband Powerline (BPL) for Internet access [5, 6, 7]. BPL and WiFi (IEEE 802.11x) are seriously considered as two other possible offerings to complement such broadband services as Digital

Subscriber Lines (DSL) and Cable TV Modems. BPL has the advantage of ease of installation with literal 'plug and play' and greater penetration inside the home. Thus the powerline, historically used for the delivery of electrical power, now also provides a high speed digital pipe to the home and a 'no new wires' communication network inside.

Multimedia In-home Networking

While HomePlug 1.0 provides acceptable data rates and performance for data communication needs in connecting multiple computers and peripherals in a LAN setting, higher data rates and more stringent QoS controls are needed to support digital multimedia communication within the home[8]. The HomePlug AV standard expected to be available in the last half of 2005, is optimized for precisely this scenario.

A single stream of High Definition Television (HDTV) may require about 25 Mbps and a typical scenario may require support for a number of simultaneous multimedia streams of voice, audio and video. Moreover multimedia applications also have latency, jitter and packet loss probability (PLP) requirements that must be met for optimal performance (see Table 1)

Application	Bandwidth (Mbps)	Latency (msec)	Jitter (nsec)	PLP (log)
HDTV	25	300	500	-10
SDTV	4	300	500	-10
DVD	6	300	500	-10
VOIP	64	10	10000	-2
Gaming	0.1	10	N/A	-6
Video conf.	1	75	10000	-6

Table 1 – Typical Multimedia QoS Requirements

Although there are several existing in-home communication technologies that appear to be capable of providing the basis for such multimedia communication, a careful examination reveals several possible deficiencies. For example, the popular IEEE 802.x suite of protocols (including the emerging IEEE 802.n standard) does not provide complete house coverage (with a single access point) at adequate data rates and reliability to provide a robust multimedia solution. Although the new Ultra-Wide band (UWB) standard will certainly have adequate bandwidth, its reach will likely be confined to a single room rather than the entire home. The recently announced standard from the Multimedia over Coax Alliance (MoCA), while possibly offering a solution for video distribution between video sources and players already connected to the exiting coaxial video cabling, fails to offer whole house coverage for other applications such as audio and VOIP, since video cabling is typically limited. Phoneline networks also have limited phone connections inside the home.

The new 200 Mbps HomePlug AV standard from the HomePlug Alliance, on the other hand, offers whole-house coverage, with an average of 44 outlets per home (in the USA). HomePlug AV will provide roughly a ten fold improvement over HomePlug 1.0, with typical TCP data rates of 50-70 Mbps, and thus it is able to support multiple simultaneous multimedia streams. Furthermore the HomePlug AV standard is

specifically designed and optimized for Audio Visual (AV) applications and will provide adaptivity to satisfy relevant QoS requirements. It should also be noted that, compared to the wireless in-home channel, the powerline communication (PLC) channel is relatively static and thus the QoS requirements are much more easily met in the more robust PLC environment.

The rest of the paper is structured as follows. Section II reviews characteristics of the powerline channel, while Section III gives an overview of the HomePlug 1.0 standard from a system perspective. Section IV provides brief descriptions of both the PHY (Physical Layer) and the MAC (Medium Access Control) protocols of the HomePlug AV specification. It describes the HomePlug AV framing structures and unique channel adaptation mechanisms. It also presents the associated network architecture that supports both Time Division Multiple Access (TDMA) as well as Carrier Sense Multiple Access (CSMA), with multiple independent overlapping networks. The paper concludes in Section V with some observations on the efficiency and performance of HomePlug AV and comments on further work in this area.

II. PLC CHANNEL CHARACTERISTICS

Multipath Channel Effects

In-building electrical wiring, designed for carrying electrical power at 50 or 60 Hz, consists of a variety of conductor types and sizes connected almost at random. The resulting terminal impedances vary both with communication signal frequency and with time as the load patterns at the consumer premises change. The net result is a multipath effect that causes delay spread (averaging a few microseconds) and deep notches (from 20 to 70 dB) at certain frequencies within the band used by PLC

communications [9]. In North America, HomePlug 1.0 uses a frequency band 4.5-20.7 MHz, while HomePlug AV uses the band from 1.8 to 30 MHz band. Regulatory constraints make frequencies above 30 MHz unattractive for PLC applications.

PLC Channel Noise Issues

In addition to the inherent fading attenuation and phase characteristics of the PLC channel, high speed communications in this channel must also mitigate a plethora of impairments and noise events which have been historically a major impediment to high speed PLC. Typical noise sources are certain types of halogen and fluorescent lamps, switching power supplies, brush motors, and dimmer switches. Furthermore, the PLC channel is subject to interference from, and without spectral masking would itself adversely impact, other users of the specified spectrum, such as citizen band and amateur radio.

Another characteristic of the PLC channel that has an impact on achievable data rates is the cyclic variation of noise with the powerline cycle. In particular, it has been found that the signal to noise ratios are much better in the vicinity of the zero crossings of the 50/60 Hz powerline cycle.

Taming the Shrew-like PLC Channel

Several specific techniques are used in HomePlug 1.0 and HomePlug AV to conquer the many hurdles posed by the PLC channel; these are described below.

Orthogonal Frequency Division Multiplexing (OFDM): OFDM is ideal for the frequency selective PLC channel since it allows the division of the available spectrum into a large number of smaller, independent flat fading channels for each of which

appropriate adaptive multi-level modulation schemes can be selected.

Programmable Spectral Masking: In order to meet regulatory constraints and to minimize mutual interference, a fixed spectral mask can be programmed such that the PLC devices do not use or cause interference in certain specified bands.

Orthogonal Channel Adaptation, Modulation and Coding: A robust, relatively low data rate scheme (*ROBO*) featuring high time and frequency redundancy, low order modulation and very powerful error coding is designed to reach almost all nodes in the PLC network. In addition, each PLC packet has a *Frame Control (FC)* segment that uses a highly reliable scheme to ensure that key parameters critical to the functioning of the PLC system, are reliably received by all nodes in the network. *Tone Maps* are used for high speed communication between a specific pair of nodes, to communicate the particular OFDM carriers and modulation and error coding schemes to be used. Tone maps are adjusted periodically based on on-going channel monitoring.

Efficient Medium Access Control Framing and ARQ: PLC communication uses a highly efficient MAC/PHY framing strategy to ensure low overheads. Furthermore channel contention, reservation and backoff mechanism are optimized to maximize throughput. Also, high speed PLC features a carefully crafted error detection and retransmission (ARQ) strategy to ensure reliable communication even in the unfavorable channel conditions.

III. HIGHLIGHTS OF HOMEPLUG 1.0

The 14 Mbps HomePlug 1.0 standard was released in 2000 by the HomePlug Powerline

Alliance to provide a PLC-based in-home LAN solution. HomePlug 1.0 stations use the well known carrier sense multiple access with collision avoidance (CSMA/CA) technique for medium sharing. This mechanism is augmented with an enhanced back-off algorithm along with priority resolution slots. The back-off algorithm enables the HomePlug 1.0 network to operate at high efficiency under varying network loads. The priority resolution slots enable four levels of strictly differentiated QoS to traffic based on priority level.

HomePlug 1.0 PHY

Orthogonal Frequency Division Multiplexing (OFDM): In HomePlug 1.0, 128 evenly spaced carriers are specified in the range 0-25 MHz. A programmable tone mask is used (in default configuration) to identify the 84 carriers that fall inside the 4.5-20.7 MHz range, among which eight are permanently masked to avoid conflict with Amateur radio bands.

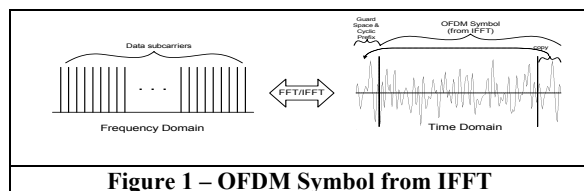
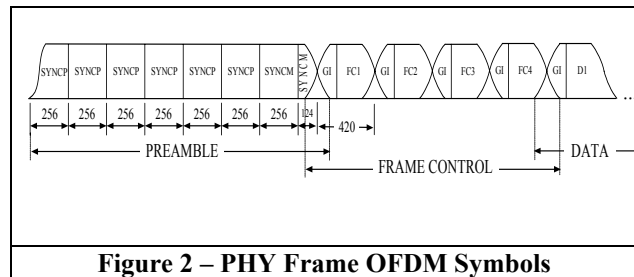


Figure 1 illustrates the generation of an OFDM symbol from the unmasked carriers via an IFFT process. Each OFDM symbol is 8.4 μ s long, with 5.12 μ s (256 samples) corresponding to the new OFDM symbol and 3.28 μ s being a cycling prefix obtained from the last 172 samples. Figure 2 shows the key elements of the PHY Frame, namely the *Preamble*, the *Frame Control (FC)* and a variable number of *Data (Payload)* OFDM symbols. Figure 3 gives further details..



Preamble: The preamble is constructed from 7.5 special OFDM symbols without cyclic prefixes, and is 38.4 μ s in duration. This frame segment is used for synchronization, automatic gain control and optionally for phase reference. The time needed to detect the preamble is the Physical Carrier Sense (PCS) interval and dictates contention slot size.

Frame Control: The Frame Control consists of 4 OFDM symbols in which all unmasked carriers are used in conjunction with a Turbo Product Code (TPC) for error correction. The four FC symbols contain 25 information bits received with high reliability. These bits, structured as shown in Figure 3(b), provide information for the correct operation of the HomePlug 1.0 protocol. (See [3] for further details).

Payload: The Payload consists of a variable number of 20- and 40-OFDM symbol blocks, protected by Reed-Solomon/Convolutional concatenated encoding. HomePlug 1.0 features fairly smooth adaptation from 1 to 14 Mbps; some 140 intermediate rates are supported by combining available modulation schemes and FEC coding rates. The effect of this smooth adaptation is seen in the stability of HomePlug 1.0 in good and bad channels when compared with the larger variations from IEEE 802.11 a and b. (See Figure 4)

Priority Resolution (PR): The two Priority Resolution symbols (see Figure 3(a)) each consist of six OFDM symbols, lasting 30.72 μ s and are used to establish four levels of

priority. The PR slots are 35.84 μs long, which takes into account the time to process the 30.72 μs symbols. Like the Frame Control, both Preamble and PR symbols must be received reliably by all nodes in the network, so all unmasked OFDM carriers are modulated and encoded in a standard way.

Figures 5 and 6 show the HomePlug 1.0 Transmitter and Receiver block diagrams. The Transmitter block shows the separate processing of Frame Control (FC) and Data bits. The data is first scrambled, then encoded, punctured, and interleaved before being mapped according to selected tone maps onto the OFDM carriers. After IFFT, a

preamble and cyclic prefix are inserted (if needed) followed by a shaping filter to effect sharp band edges. At the receiver a synchronization block detects the presence of a preamble signal, and the subsequent frame control and data symbols undergo receive side processing to de-modulate data, and to de-code the bitstreams of interest.

For ease of implementation, cost and other reasons, HomePlug 1.0 uses differential modulation with only DBPSK and DQPSK schemes. In addition all carriers used in each OFDM symbol have the same modulation scheme.

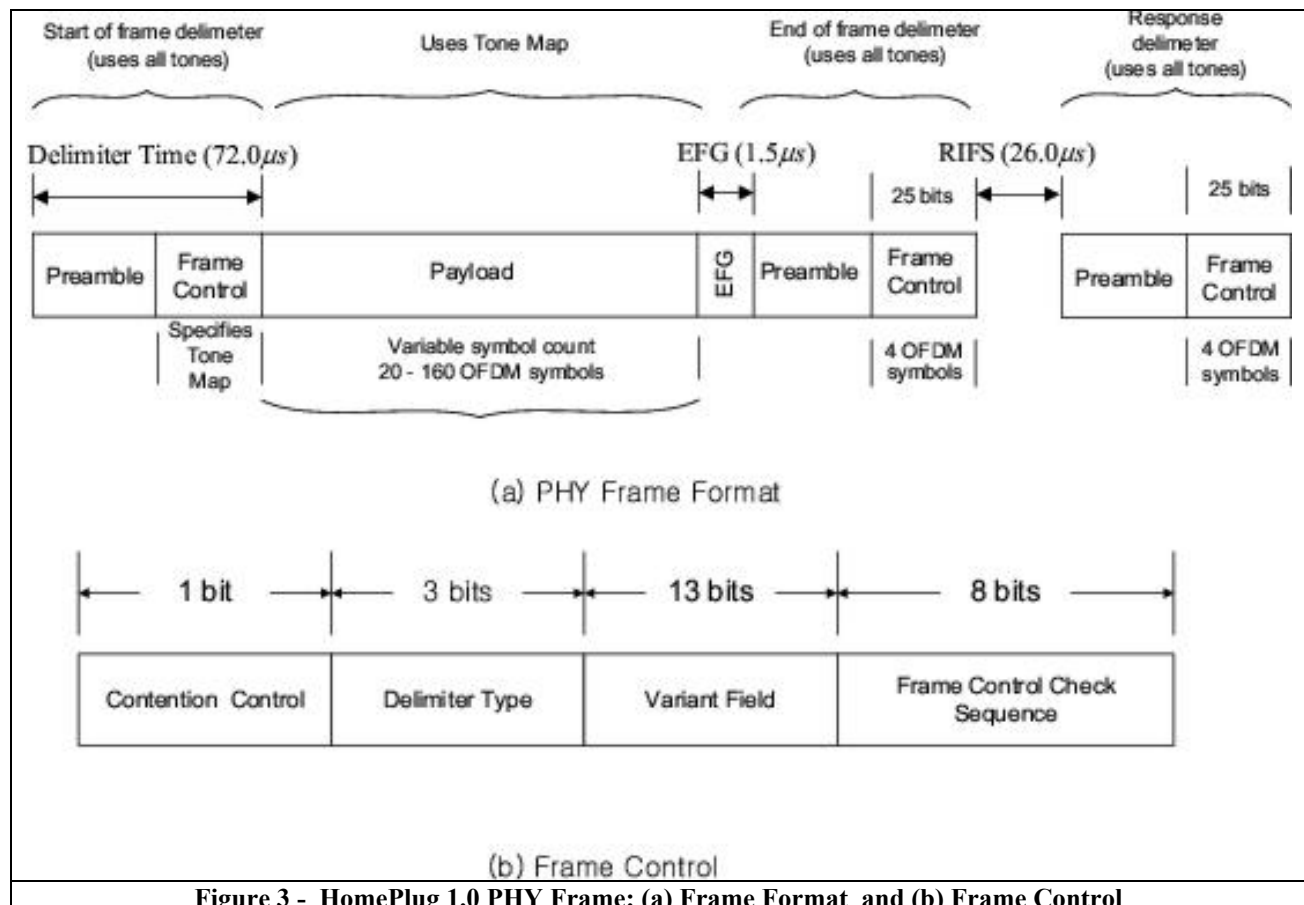


Figure 3 - HomePlug 1.0 PHY Frame: (a) Frame Format and (b) Frame Control

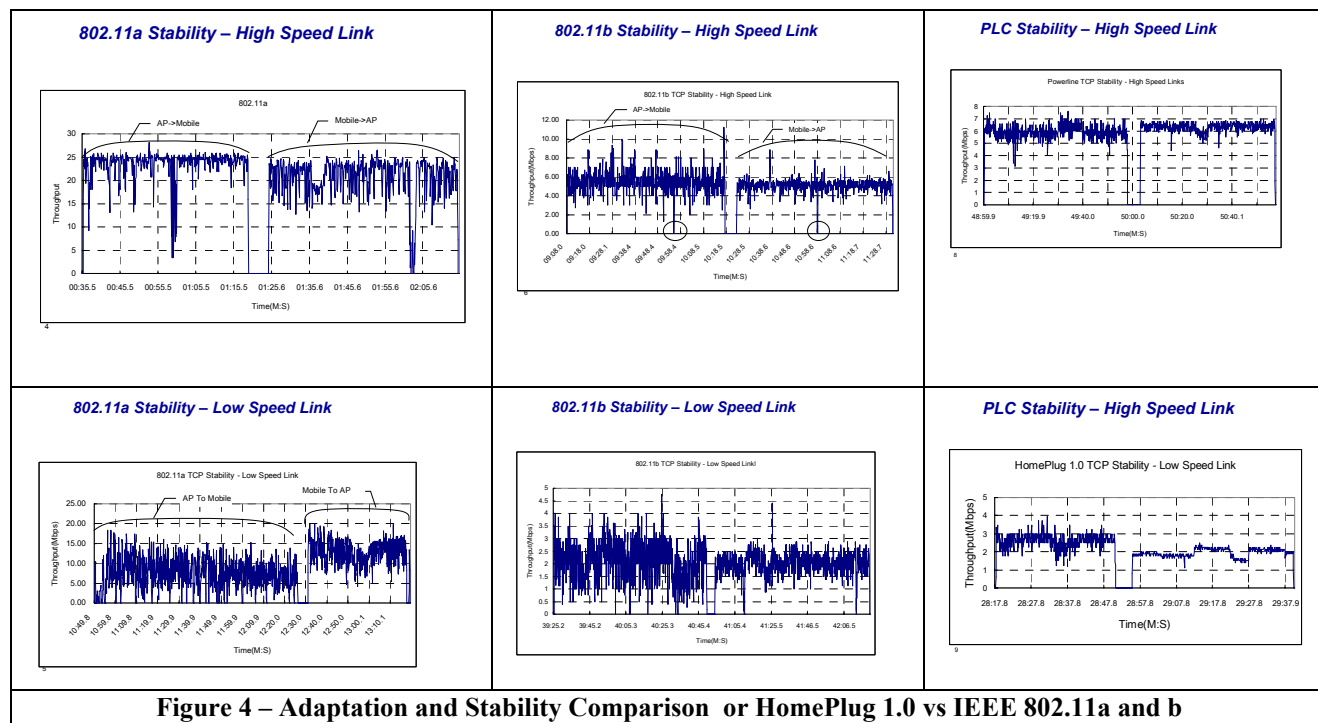


Figure 4 – Adaptation and Stability Comparison of HomePlug 1.0 vs IEEE 802.11a and b

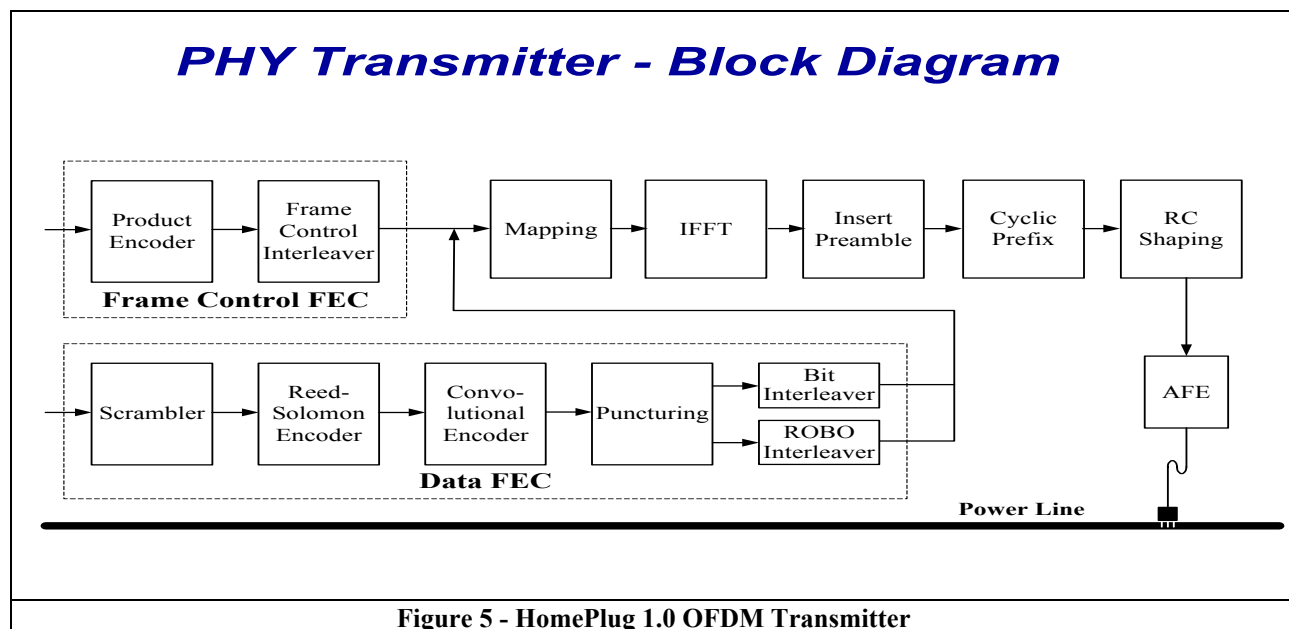
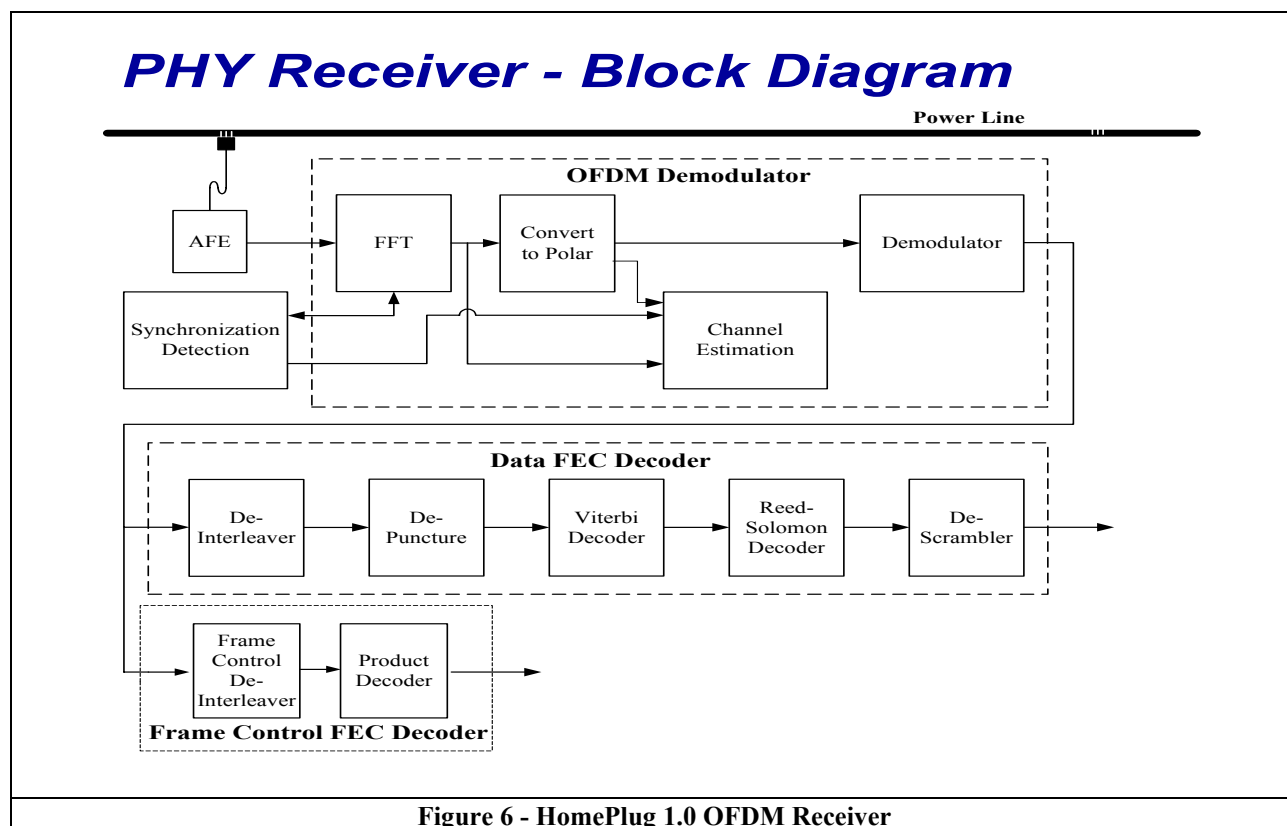


Figure 5 - HomePlug 1.0 OFDM Transmitter



HomePlug 1.0 Medium Access Control

HomePlug 1.0 uses Carrier Sense Multiple Access with Collision Avoidance (CSMA/CA). Physical Carrier Sense (PCS) is complemented by Virtual Carrier Sense (VCS) information contained in the Frame Control Field indicating whether other stations can contend for the medium or not. Figure 3(a) shows the basic frame structures and timing involved in Medium Access Control. The payload is prepended with a delimiter constructed from a Preamble and Frame Control as described earlier. After a period denoted End of Frame Gap (EFG) of 1.5 μ s, an End-of-Frame (EOF) delimiter is added by repeating the Preamble and Frame Control, thus increasing the likelihood that all nodes will be synchronized and correctly receive the important information contained in the Frame Control fields.

All nodes wait for a period of Response InterFrame Spacing (RIFS) of 26 μ s for the

response to be sent in the form of a response delimiter consisting of the preamble sequence and the Frame Control symbols. The Frame Control fields shown in Figure 3(b) contain VCS, priority, and acknowledgement information, needed for the proper operation of the protocol.

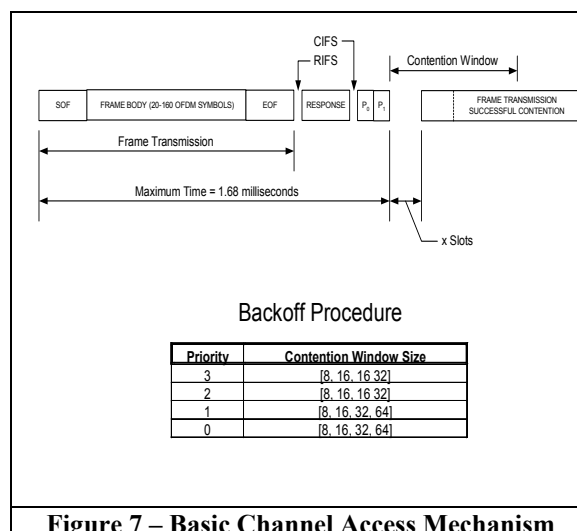
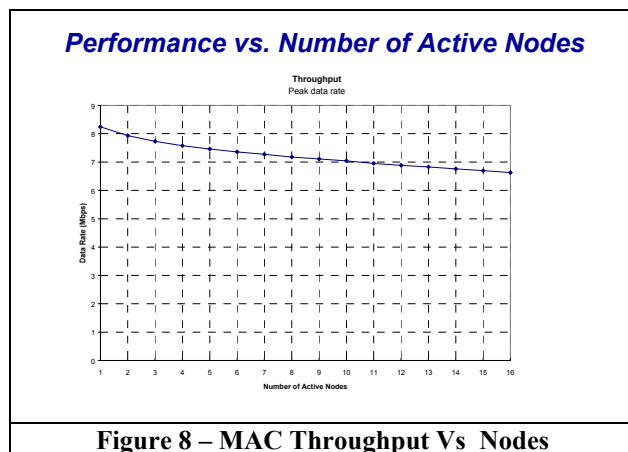


Figure 7 shows how channel contention proceeds at the end of the response. If the Contention Control (CC) bit is set in the Frame Control of the response, then the node presently sending data at a certain priority will continue to attempt to send data ("bursting" - only Priority CA2 and CA3 nodes are allowed to do this), but could be preempted by higher priority nodes asserting their priority in PR0 and PR1. Priority slot PR0 begins after an interval of Contention Interframe Spacing (CIFS) from the end of the response delimiter. Nodes seeking access to the channel must first assert priority CA0-CA3 in PR0 and PR1.

Nodes of the highest winning priorities will then contend for the medium in the contention slots and the node that wins the contention begins to transmit data. Colliding and losing nodes will choose new backoff values from the backoff window for that priority according to the backoff schedule shown in Figure 7. Note that unlike the 802.11x backoff procedure, both colliding and losing nodes in the contention procedure may choose new backoff values, depending on a new variable called *the deferment counter*, which checks how many times a particular node has lost contention.



Security and Key Management

HomePlug 1.0 uses a password-based cryptography standard (PBCS) for key management to effect cryptographic isolation of logical networks. All stations in a logical network share the same Data Encryption Standard (DES) key, called a Network Encryption Key (NEK). Encryption is enabled by default and cannot be disabled, but for proper protection, the user must select a unique network password.

HomePlug 1.0 Performance

Simulations and measurement show that HomePlug 1.0 provides typical throughputs of 5-7 Mbps (TCP), Full house coverage in 99% of the homes tested was observed with a data rate of at least 1.5 Mbps. Figure 8 shows how the theoretical MAC throughput for HomePlug 1.0 varies with the number of nodes.

Further details of the operation of the HomePlug 1.0 protocol and the functions of the various fields, such as Frame Control, are contained in [3].

IV. OVERVIEW OF HOMEPLUG AV

HomePlug AV Bandwidth

HomePlug AV provides an order of magnitude throughput improvement over HomePlug 1.0, while also addressing key QoS issues. The bandwidth used has been extended and subcarrier spacing reduced in AV. Whereas HomePlug 1.0 uses 4.5 to 20.7 MHz quantized into 84 subcarriers, AV operates with 1155 carriers over 1.8 to 30 MHz. While Homplug 1.0 in its default configuration uses 76 active carriers in its bandwidth of operation, Homeplug AV uses 917 in its default mode.

HomePlug AV OFDM Symbol

Similar to the HomePlug 1.0 standard, Orthogonal Frequency Division Multiplexing (OFDM) is used for HomePlug AV. However, various OFDM system parameters have been updated in order to maximize spectral mask flexibility and increase system throughput. Figure 8 shows the structure of the HomePlug AV symbol and Table 3 gives the values of the key PHY parameters.

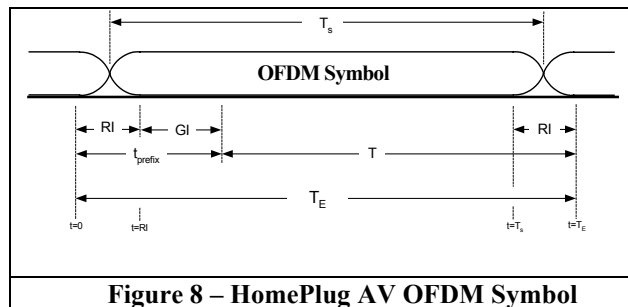


Figure 8 – HomePlug AV OFDM Symbol

The OFDM symbol's IFFT interval time in HomePlug AV is approximately eight times that of HomePlug 1.0. One advantage of this is that, in the basic configuration, (5.56μs or 7.56μs guard interval) the overhead due to the guard interval, used to mitigate intersymbol interference (ISI), is much less in HomePlug AV. Likewise, when the system encounters a channel where the delay spread is larger than the guard interval, subcarrier SNRs are not impacted as greatly due to the fact that the percentage of the IFFT interval affected is less.

Another advantage of the longer symbol time is that the OFDM symbols can be (and are) shaped and overlapped in such a way that deep frequency notches can be created simply by turning carriers off, whereas HomePlug1.0 required, either turning off a

Table 3 – HomePlug AV OFDM Symbol Characteristics

Symbol	Description	Time Samples	Time (μs)
T	IFFT Interval	3072	40.96
t_{prefix}	Cyclic Prefix Interval	RI+GI	4.96+GI
T_E	Extended Symbol Interval ($T + t_{\text{prefix}}$)	$T + t_{\text{prefix}}$	45.92+GI
RI	Rolloff Interval	372	4.96
T_s	Symbol Period	3072+GI	40.96+GI
GI_{FC}	Frame Control Guard Interval	1374	18.32
GI	Data Symbol Guard Interval, generically	417, 567, 3534	5.56, 7.56, 47.19
GI_{417}	Guard Interval, length=417 samples	417	5.56
GI_{567}	Guard Interval, length=567 samples	567	7.56
GI_{3534}	Guard Interval, length=3534 samples	3534	47.19

large number of carriers both in and around the desired notched band, or additional filtering. Figure 9 details the carrier power rolloff for the three guard intervals. Though it varies with guard interval, it can be seen that if all carriers within approximately 115kHz of a desired notched band are turned off, the energy will be at least 30dB down in the notched band. Figure 10 shows the deep notching achieved in HomePlug AV.

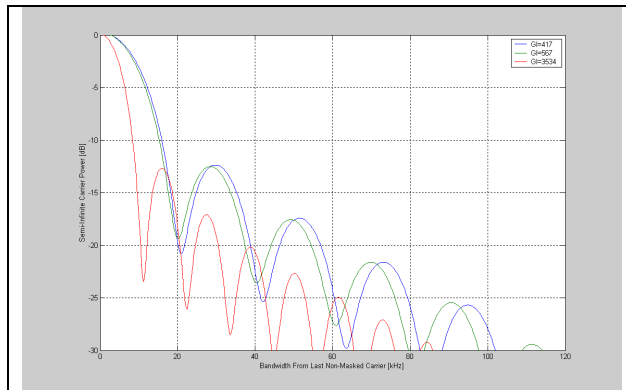


Figure 9 – Notching by turning Off Carriers

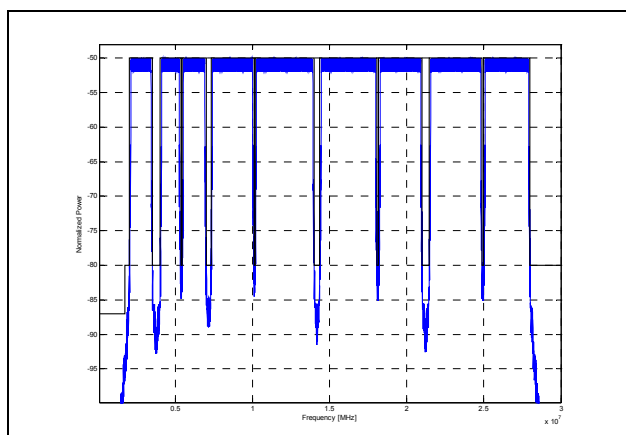


Figure 10 – Deep Notching in AV

HomePlug AV Carrier Modulation

Carrier modulation has been improved in HomePlug AV to maximize channel throughput. HomePlug 1.0's differential modulation has been replaced in HomePlug AV with coherent modulation – yielding higher carrier SNRs for a given signal power. Second, whereas HomePlug 1.0 used only DBPSK or DQPSK modulations,

individual HomePlug AV carriers can be modulated with BPSK, QPSK, 8-QAM, 16-QAM, 64-QAM, 256-QAM, or 1024-QAM. This allows the system to take full advantage of all possible ranges of SNRs that a particular subcarrier could encounter. Finally, in contrast to HomePlug 1.0 that does not mix modulation types across carriers, HomePlug AV fully supports bit-loading. A mix of modulations is tailored for each channel such that each carrier communicates with the fastest modulation that the carrier's SNR can support.

HomePlug AV FEC

Forward error correction (FEC) has also been improved in HomePlug AV. Whereas HomePlug 1.0 uses a concatenated code, HomePlug AV uses a state-of-the-art turbo convolutional code, allowing greater throughput for a given channel SNR, a gain equivalent to about 2.5 dB. While HomePlug 1.0 had a single ROBUst mOdulation (ROBO) scheme, HomePlug AV features several additional robust modes of operation in which a repetition code is applied as an outer code to the turbo code for broadcast or for use in harsh channel conditions.

HomePlug AV and 1.0 Coexistence

The HomePlug AV technology was designed to be able to coexist with HomePlug 1.0 nodes in a given network. HomePlug AV has the ability to send delimiters recognizable by HomePlug 1.0 nodes in order to communicate protocol information regarding channel access and contention.

The major elements of the HomePlug AV transmitter and receiver are shown in Figure 11. Note that in the transmitter, HomePlug 1.0 Frame Control, HomePlug AV Frame Control and HomePlug AV packet body are generated separately, and are similarly decoded independently at the receiver.

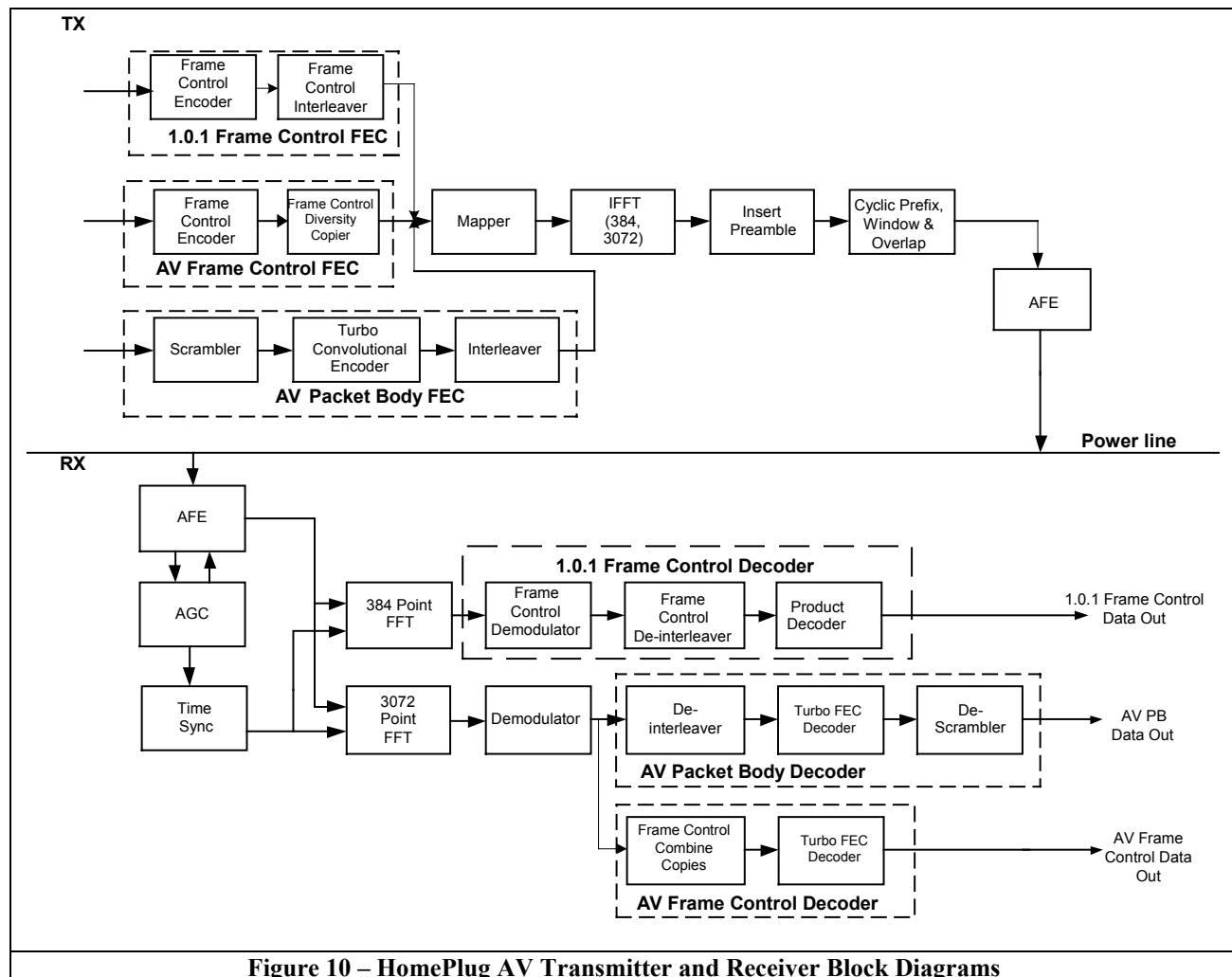


Figure 10 – HomePlug AV Transmitter and Receiver Block Diagrams

HomePlug AV Medium Access

In HomePlug AV, medium access is primarily through Time Division Multiple Access (TDMA), with CSMA/CA available for bursty applications. In each network, a Central Coordinator (CCo) transmits a beacon frame that contains schedule information for the other stations. Stations that source steady streams request time allocations from the CCo, and transmit in the assigned regions. This avoids the overhead of contention and collision present in CSMA/CA.

Framing and Segmentation

In HomePlug 1.0, relatively low PHY rates made it reasonable to transmit a single

incoming host packet in one or more MPDUs (MAC Protocol Data Units). The order of magnitude increase in PHY rates achieved by HomePlug AV make this approach very inefficient, so incoming host packets are aggregated into a stream of MAC frames, with a total of six bytes of header and ICV (integrity check value) per MAC frame. The MAC frame stream is then segmented into fixed length blocks called PHY Blocks (PBs) that are independently encrypted and corrected. One or more PBs are sent in each MPDU, with each PB carrying a four-byte header to allow correct reassembly.

SR-ARQ Error Control

HomePlug AV employs Selective Repeat Automatic Retransmission Request (SR-ARQ). Each PB has its own 32-bit Cyclic Redundancy Check (CRC) to detect errors. The receiver responds with a Selective Acknowledgement (SACK) that pinpoints the PBs requiring retransmission. Only the damaged PBs are retransmitted, and these may be combined in a new MPDU with newer PBs that are being sent for the first time. This approach allows full MPDUs to be sent almost all the time, so that the fixed delimiter overhead remains small relative to the total transmission time.

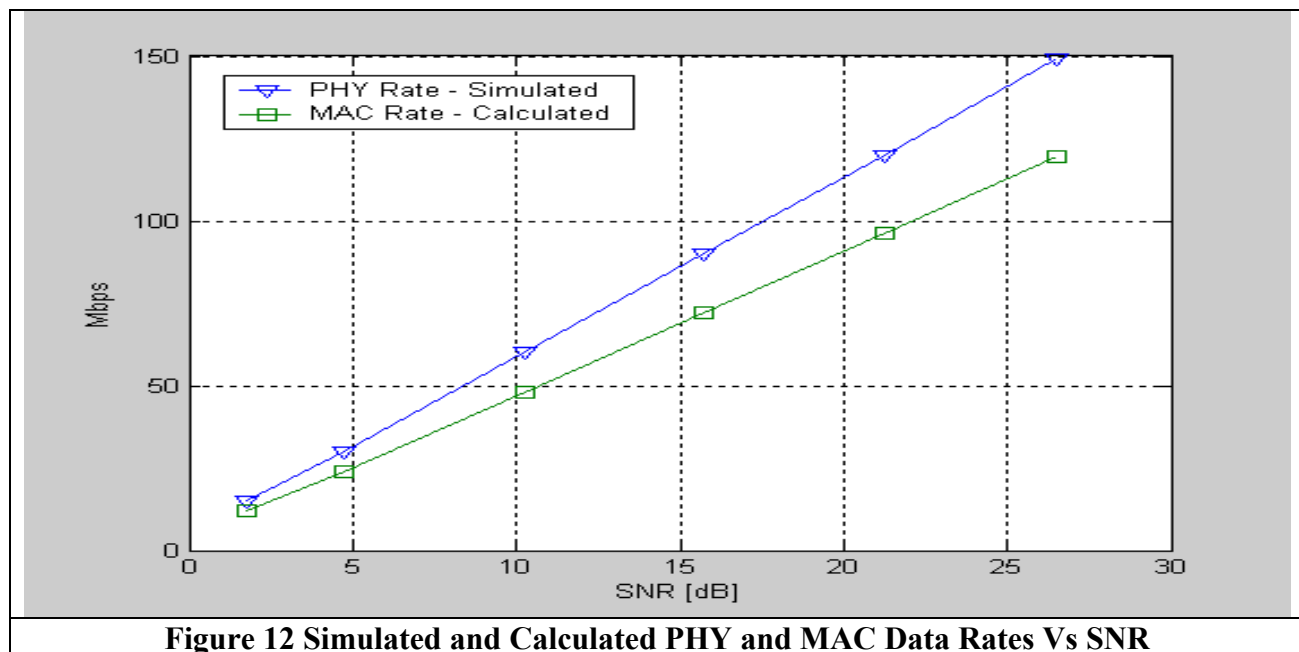
Security and Key Management

While HomePlug 1.0 uses 56-bit DES encryption, HomePlug AV uses 128-bit AES. Both use Cipher Block Chaining (CBC) to increase randomness in similar

transmissions. The Initialization Vector (IV) is transmitted explicitly in HomePlug 1.0, whereas in HomePlug AV, it is derived from frame information.

QoS in HomePlug AV

To support desired delay, packet loss tolerance, and jitter, HomePlug AV takes several measures. As explained above, access for steady streams (such as multimedia applications generate), is carefully scheduled using TDMA. Allocated times reflect the latency requirements, and provide sufficient time for retransmissions as needed to meet the PLT requirements of the stream. Jitter is managed by timestamping incoming data units with their target delivery time.. Stations execute a time synchronization method to remain in tight synchronism so that the jitter remains below 500 ns.



V. COMMENTS AND CONCLUSIONS

HomePlug AV has made many significant improvements over the already successful HomePlug 1.0 protocol. It is much more

efficient and provides stringent QoS guarantees that are impossible to meet in HomePlug 1.0.

HomePlug AV Performance

The improved design of both PHY and MAC in HomePlug AV render it tremendously efficient. At the PHY level, the data rates achieved are very near the information theoretic limits.

MAC framing overhead is minimized and the error correction and retransmission scheme provides an excellent combination of reliability and efficiency. Typical MAC efficiencies are projected to be in the 80% range, depending on the nature of the application and the PHY rate.

HomePlug AV is capable of complete house coverage and will support multiple streams of high and standard definition television, and stereophonic hi-fi music, while still supporting high speed data applications. Figure 11 shows the HomePlug AV PHY Rate and MAC throughput, as a function of signal to noise ratio. No other technology can provide such data rates with whole-house coverage and thus HomePlug AV is expected to provide synergistic solutions for home entertainment equipment manufacturers and content providers.

Broadband Powerline access (BBL) is also certain to benefit from the emerging set of high speed PLC chips, with symmetric access speeds of 20-40 Mbps or more expected to the the home. Due regard is being given to designing emerging standards so BPL and PLC LANs can co-exist. In this regard it is worthwhile to note the recent formation of the IEEE Technical Committee on PLC which is actively promoting PLC and BPL research, and forming relationships with traditional academic and professional communications organizations.

VI. REFERENCES

1. Yu-ju Lin, Haniph A. Latchman, Minkyu Lee and Srinivas Katar, "*Power line Communication Network Infrastructure For Smart Homes*", **IEEE Wireless Commu-nications**, Volume 9, Issue 6, Pages: 104-111, December, 2002.
2. HomePlug Powerline Alliance, <http://www.homeplug.org> [February 10, 2005]
3. M.K. Lee, , R. Newman, H.A. Latchman, S. Katar, and L. Yonge, "*HomePlug 1.0 Powerline Communication LANs –Protocol Description and Comparative Performance Results*", accepted for publication in the Special Issue of **the International Journal on Communication Systems on Powerline Communications**, pages 447-473, May, 2003
4. Y. Lin, H. Latchman, S. Katar, and M.K. Lee, "*Theoretical and Field Performance Comparison between 802.11 Wireless and Powerline HomePlug 1.0 Protocols*" accepted for publication in the **IEEE Communications Magazine with Focus Theme on Powerline Local Area Networks**, pages 54-63, May, 2003.
- Current Technologies [February, 2005] <http://www.currenttechnologies.com>
5. Amperion Corporation, February 2005, <http://www.amperion.com/>
6. Ambient Corporation February 2005, <http://www.ambientcorp.com>
7. Baowei Ji, Archana Rao, Minkyu Lee, Haniph A. Latchman, Srinivas Katar, "*Multimedia in Home Networking*" - **CITSA 2004 / ISAS 2004 Proceedings Volume 1, Network Technologies**, Page Nos: 397-404

8. Barnes, J.S., "*A Physical Multipath Model For Powerline Channels at High Frequencies*", **Proceedings of the International Symposium on Powerline Communication and Its Applications**, 1998, 76-89.
9. Minkyu Lee, Haniph A. Latchman, Richard E. Newman, Srinivas Katar, and Larry Yonge, "*Field performance comparison of IEEE 802.11b and HomePlug 1.0*", **27th Annual IEEE Conference on Local Computer Networks**, Pages: 598-599, 2002.

HOME NETWORKING ON COAX FOR VIDEO AND MULTIMEDIA

Ladd Wardani
Entropic Communications

Abstract

Home networking of multimedia, and particularly of video, is covered, including usage models, system requirements, installation and maintenance, security, and comparison of the various in-home mediums. Field characterization of the in-home coax plant is described, followed by home field testing data of the technology developed by Entropic and field tested by MoCATM.

INTRODUCTION

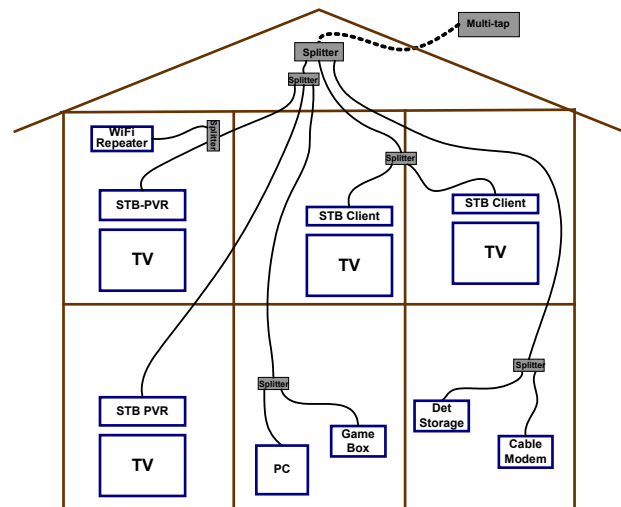
The ability to home network digital entertainment, including multiple video, audio and data streams has received significant attention and effort in recent years. This effort is driven by the desire to have content from DVRs, audio devices, PCs, and broadband (including FTTH) available anytime and anywhere throughout the home. With the major new application being sharing video to all the home's displays, many product developers and service providers have realized that the "missing-link" to this capability is the availability of a ubiquitous, highly reliable and high throughput home network.

While consumers will accept lower quality or degraded video viewing on mobile and hand held devices, even occasional glitching, blocking or "buffering please wait" will not be tolerated for entertainment with standard home video devices such as TVs, flat panel displays, VCRs, DVRs, STBs, and media centers. This paper focuses on home networking of high quality video, defined as

HDTV, SDTV and DVD without degradation, and is consistent with current operators' broadcast reliability.

USAGE MODEL

The following figure shows the home usage model with triple play services (voice, video, data). This home is using coax home networking, however the desired usage model is independent of home networking medium.



Usage Model Content

The home network must support multiple and simultaneous HDTV, SDTV, audio, data, VoIP, gaming, and other multimedia usages both from the broadcast network and from an in-home DVR or storage device. Video content will continue to be over MPEG2 as well as advanced coding (MPEG4, H.264, WM9, etc.).

Usage Model Data Flow

Each room and device may be either, or both, a source or sink of content both to and from multiple simultaneous devices. Consumers may add or move devices from room to room, changing where sources and sinks connect. Flow within a home may traverse the home network twice (double hop), such as when video is sent from a point of ingress to a DVR, and then from the DVR to a display.

Usage Model Installation

Devices may be provided from either retail or the service provider. While service providers may choose to professionally install these systems initially or perpetually, they do not want to choose a home networking technology that precludes a retail self-installation model.

REQUIREMENTS

The requirements on the home network are numerous. However, the key requirements are reliability and ubiquity, without which deployment is not feasible. There does not exist any reasonable home networking solution that provides 100% immediate penetration and ubiquity. Even 100 Mbps Ethernet over cat-5 has been shown not to be a 100% solution with issues including that rate negotiation problems can drop or bounce networks to 10 Mbps, many of the devices being deployed in homes can not in reality support more than 40 to 60 Mbps of real video, consumers can add hubs and collision domain devices to the network, and most devices lack QoS. The home networking solution that can and will be deployed is one that is:

- A >95% solution, with reasonable remediation for the remaining <5%.

This is the key criterion by which the home network must be judged .

With multiple HDTV and SDTV streams requiring the vast majority of data rate and the best packet error rate (PER), plus security, a home network that can support multiple HDTV streams meets the main requirements. The addition of low jitter and low delay for voice and gaming then covers it. Here, then, are the essential requirements:

- Coexist with existing services
- Medium collocated with target devices
- No changes to wiring or splitters or other medium specific devices
- Full mesh, peer-to-peer network
- Data rate net (MAC) > 60 Mbps
 - 100 Mbps preferred
- PER < 1e-6, BER 1e-9
- Delay < 20 msec
 - < 10 msec preferred
- Jitter < 1 msec
- Privacy from neighboring homes
- No degradation due to neighbor's home networking or general appliances
- No degradation due to other in-home networking products or general appliances
- No retransmissions above MAC layer
- Plug n' play on its medium
- Carry key protocols (Ethernet, 1394)
- Support various DRM
- Meet consumer price points
- Support multiple independent networks
- Service provider may open the home network to CE devices or keep it closed
- Futureproof to higher data rates

Key Requirement Drivers

Data Rate Requirements

An often overlooked but crucial point is that the home network must support the peak data rate per stream, and that streams can peak simultaneously. The vast majority of broadcast content in the USA today is MPEG2. Measurements of streams from the top 4 USA service providers showed that SDTV MPEG2 broadcasts are variable bit rate with average data rates between 2 and 3 Mbps, and peak rates around 9 Mbps. VOD SDTV streams tend to be constant bit rate between 4 and 5 Mbps.

HDTV MPEG2 streams can carry requirements from the content providers that a minimum of around 12 Mbps is allocated. Measurements of programs in San Diego showed that non-ATSC HDTV MPEG2 streams carry average data rates between 10 and 18 Mbps and peak rates between 13 and 19 Mbps.

Advanced coding schemes may reduce these average bit rates by around a factor of 2; however the peak rate does not reduce by as much as the average rate.

Fast forward and reverse, and other trick modes, can increase the peak data rate by a factor of 3 or more if continuous looking video is desired during the trick mode. If a decimated, fast slide show looking trick mode is acceptable, then data rates stay comparable.

Service providers will still need to make terrestrial or network ATSC content available to their subscribers. So even if an advanced codec brings the service provider's data rate down for HDTV, the subscriber may sometimes home network ATSC. Service providers must allocate the home networking

data rate for ATSC streams at 19 Mbps regardless of their own HDTV data rates.

Table 1. Data Rates in Mbps			
	Ave	Peak	Trick Modes
SDTV	1 – 3	3 – 9	3 – 20+
HDTV	6 – 18	9 – 19	9 – 40+
ATSC	19	19	19 – 40+
Double Hop	x2	x2	x2

Then, for example, a usage scenario with 1 HDTV stream, 1 ATSC stream, 2 SDTV double hops, and broadband at 10 Mbps requires allocation of 50 to 84 Mbps for play, and 50 to 148 Mbps if all streams did trick mode at the same time. A usage scenario with 3 ATSC streams and broadband at 3 Mbps requires allocation of 60 Mbps for play, and 60 to 123 Mbps for trick modes.

None of these numbers include another typically > 10% protocol and management overhead from the application. Additional data rate of 30% or more of a stream can also be needed to fill IPSTB buffers at channel changes in order to reduce channel change times.

Reliability Requirements

Digital cable programming is delivered with threshold PER of below 1e-6. The home network should have similar or better performance so as not to degrade viewing. When supporting UDP and 1394 protocols there is no request and retransmission of packets above the Link/MAC layer, so that this PER must be constantly maintained at the MAC layer by the home network. A home network with potential collision access will have a very difficult time achieving 1e-6 PER with any significant loading of its data rate. A fully coordinated MAC, collision free, is practically a necessity for currently

reasonable home networks to achieve this PER for the desired data rates.

Reliability from Appliance and Neighbor Interference Requirements

The home network should not degrade due to other networking devices or typical electronic devices and appliances such as cordless phones, wireless laptops, wireless hot spots, vacuum cleaners, power drills, hair dryers and microwave ovens. This must be true for both the subscriber's other devices as well as the neighbor's devices. If reliability were subject to other devices and/or the neighbor, then not only would there be service interruptions but they would not be predictable or easily diagnosable.

Imagine a service provider's installer completing a successful installation of several home networked devices, rolling away in his truck, only to soon return when the subscriber experienced interruptions due to an appliance that when the installer returned, was no longer on, and there was no longer any interruption. Such a situation is not reliable, maintainable, or viable.

Quality of Service Requirements

In addition to a $PER < 1e-6$, the home network must deliver delay and jitter consistent with the services it carries. Gamers claim to be able to sense delays between 10 and 20 msec. Voice applications would like to see the delay budget allocated to home networking to be 10 msec or less. Jitter must be smoothed out in buffers at the transmitter and receiver. Such buffers must be a fraction of the total allowed delay.

Bandwidth Management Requirements

When the aggregate content's data rate exceeds, on a peak or average basis, the

home networking data rate, then something must be delayed and/or dropped. Prioritization, for example using 802.1p, must be used to delay or drop lower priority traffic (typically data traffic). Video packets can not generally be delayed or dropped and must either have their data rate fully supported at $PER 1e-6$ or better, or else the video will be degraded... If there is not sufficient home networking aggregate data rate remaining to support a video stream or other traffic that can not accept dropped packets beyond a $1e-6$ rate, then a bandwidth manager must inform the home owner that the traffic can not be supported and provide options for stopping other streams or services.

In practical usage then, where data rate is taken up by high speed videos, the important elements to ensure a quality experience are:

- very high aggregate data rate
- guaranteed packet delivery at $PER < 1e-6$
- a bandwidth manager and user interface

Installation Requirements

No new wires or changes of any sort to the in-home infrastructure should be the $\geq 95\%$ rule.

A PC must not be required for self-installation. The target subscribers are TV viewers and not necessarily PC users. Even PC households do not expect to need a PC to install or operate their video systems.

MEDIUMS AND TECHNOLOGIES

The mediums in the home are wireless, phone line, powerline and coax. Each medium has innate characteristics, advantages and disadvantages. Each of these

mediums has had one or more home networking technologies developed for it.

Selecting the Right Medium

For triple play, the medium selection is dominated by video, since all of the mediums can reasonably distribute the 1 Mbps received through broadband service today. The broadband service home networking will usually be wireless for mobility and laptop use. Video home networking, however, has very different requirements outlined above. Using these video requirements, the home mediums are evaluated below.

Wireless for Video Home Networking?

Numerous wireless solutions have been proposed and implemented. Despite its inherent attractiveness, and significant technology advances, no reasonable wireless solution has yet achieved the requirements for “whole-house” home networked video.

Furthermore, service providers can not subject their revenues to the potential of propagation and interference problems from existing and future other ever-expanding ISM unlicensed band wireless products and services. There can be no guarantee that services will be maintainable in an unlicensed band due to non-controllable and non-predictable interference. As a result, there exists no premium video service in an unlicensed wireless band.

Unfortunately, no technology can remove the “un” in unlicensed. Only the FCC can. With requirements of 100 Mbps, and a frequency band that should support going through walls, the economics do not make any sense for a licensed wireless home networking band to be created.

Real life tests of wireless “whole house” home networking shows that even in benign environment and optimistic scenarios, existing wireless solutions cannot come close to meeting the minimum data rates described above. Practical rates in homes are below 20% of the minimum required rates.

The wireless medium enables very easy pirating of a single subscription since it is very difficult to control wireless network reach. A neighbor’s living room is often much closer to an access point than the furthest bedroom. Installers have a much clearer understanding of existing wired networks and will require a learning curve to service wireless networks. The remedies for a whole wireless network can be fairly complex.

Thus, whole home wireless home networking of video does not provide close to a 95% solution and does not provide a reasonable remediation solution. At best, wireless home networking of video will some day be an in-room solution, with a wired backbone for whole house coverage.

Phone Line For Video Home Networking?

Real life experience and measurements of current phone line home networking devices is and has been available for years and shows that ubiquitous whole-house coverage is a fraction of the required data rates described above.

Furthermore, in many homes, there is no phone connection in the proximity of the television sets, especially in the primary family room location.

Powerline for Video Home Networking?

Powerline connectivity seems quite attractive since it provides a wide home

coverage and is available at any location where a powered CE device is present. Powerline home networking has been used traditionally for various lighting and other low data rate control applications. For the past several years higher data rate solutions have been available. They have shown the limitations of the medium, since home testing showed that the 95% outlet can only be relied on to provide less than 2 Mbps, even though the particular technology was capable of much higher data rates.

Powerline has a similar situation as the unlicensed wireless band, but instead with effectively “unlicensed jammers”. Appliances such as vacuum cleaners, drills, hair dryers, light switch dimmers, and power supplies in many devices output noise onto the powerline that overwhelms powerline communications. These appliances can be turned on and off at any time and in different locations in the home or in the neighbors’ homes.

As with wireless, there can be no guarantee that services will be maintainable since powerline is subject to largely varying, non-controllable and non-predictable interference. This variation prevents a viable installation and service. Remediation of powerline would require a certified electrician, which cable operator installers are not.

Typically 6 to 8 homes in the USA share a transformer and therefore degrade each others performance. Powerline only has one frequency channel that all homes use. Isolating homes requires a certified electrician and expensive filter isolators.

Most all STB, PC and other device manuals recommend using a filtered power surge protector. A significant percentage of surge protectors greatly attenuate signals in

the 4 to 30 Mhz band, and block or impair powerline communications.

Thus, powerline home networking of video does not provide close to a 95% solution and does not provide a reasonable remediation solution for the problematic percentage.

Structured Wiring For Video Home Networking?

An attractive medium for multimedia home networking is the Cat-5 structured wiring widely used for business local area networks. Cat-5 (or Cat-6) supports a reliable, repeatable, viable 100 Mbps or even 1000 Mbps. Many of the new homes built today, provide Cat-5 wiring with connectors which are mostly co-located with the phone connectors.

The main drawback of the structured wiring is its relatively low residential penetration. Critical mass penetration into the tens of percentage of US homes is not expected for many years to come due to the costs of outfitting existing residences.

Other drawbacks are that the Cat-5 locations are not always near the home’s television locations, and that it is easy for the consumer to add data products, hubs and other devices that can degrade performance of the network.

Coaxial Wiring For Video Home Networking?

Coax and Cat-5 wiring stand alone as the reliable, repeatable, high bandwidth mediums in the home. However, coax is in around 100 times as many whole homes as Cat-5, and is located at over 250 million video devices in the USA, including, TVs, flat panel displays, DVRs, STBs, Media Centers, DVDs, VCRs

and Cable Modems. Almost anywhere video is being watched today in USA homes it is via coax. Coax is already validated and used today to carry many gigabits per second of video, audio and data in the 5 to 860 Mhz band, as well as the 950-2150 Mhz band for satellite operators.

Coax is a shielded medium and is not subject to interference or changes due to appliances in the home or neighbors' homes, and does not have consumer adoption issues since all consumers expect to get their video over coax. Coax is a contained medium that can not be easily shared with the neighbor for piracy, and can be physically isolated from all neighbors with reasonable effort. Usually homes can not hear each others signals, and total physical isolation between homes can be accomplished simply by putting a small \$1 filter at the POE or multita. Coax has the bandwidth to support greater than 8 non-overlapping frequency channels so that homes do not degrade each other even without physical isolation.

Coax supports sufficient bandwidth above 860 Mhz to enable multiple 100 Mbps frequency channels to coexist, while phone line, powerline and the 2.4 Ghz ISM wireless band can only hope to support one such channel best case.

Coax is the only medium that simultaneously satisfies all of the following requirements:

- Exists at > 250M home video locations
- Ubiquitously supports whole home data rates > 100 Mbps
- Capacity of many gigabits per second for future higher bandwidth implementations
- Stable. Not a time-varying quality link

- No degradation due to neighbor's home networking or general appliances
- No degradation due to other in-home networking products or general appliances
- Supports multiple independent, non-overlapping networks
- Easily physically isolated from all other homes if desired
- Does not have a slew of legacy and data home networking products that can create reliability issues for service provider's video home networking

Shared Mediums

When a medium is a "shared medium," there are two concerns that must be dealt with:

- Preventing the neighbor from snooping content or pirating service
- Preventing neighbors from interfering with each other and degrading each others service

If neighbors' mediums are physically isolated from each other this solves all issues. Otherwise, encryption is relied upon to prevent snooping or pirating of service. However, interference is only reasonably solved by creating orthogonality between neighbors on different non-overlapping frequency channels.

Wireless, powerline and coax can all be shared mediums in that the neighbor can potentially demodulate another neighbor's signal.

Powerline in the USA is connected via the transformer to typically around 6 to 8 homes, whereas in Europe it can be to more than 100 homes. Powerline communication is

restricted to low frequency operation (roughly 4 to 40 Mhz), due to attenuation, and can only support one high speed channel at best. Physically isolating homes on powerline requires a certified electrician to install expensive filters at the transformer.

Wireless can leak to many neighbors, or be purposely pointed to a neighbor with a directional antenna. The 2.4 Ghz ISM band only supports one whole home, high speed channel at best. Physically isolating home for wireless is not possible.

Coax between homes is isolated by drop cables and multitaps. Multitaps give varying amounts of isolation between tap ports that varies with frequency, and can be insufficient to fully isolate homes in the 860 – 2000 Mhz band. Multitaps come in 2, 4 and 8 tap versions. The majority are 2 and 4 tap

multitaps. Coax can support reliable 100 Mbps in a 50 Mhz bandwidth or less, enabling more than 8 channels to exist above 860 Mhz.

Physically isolating homes on coax requires simple 860 Mhz low pass filters be installed between the multitap and the POE. However, since there are more channels supported than shared homes on a multitap, each home can be operated on its own frequency and eliminate the need for physical isolation filters until the neighborhood is heavily penetrated with multiple home networking channels per home and uses more channels than are available within the total bandwidth. Thus, in combination with encryption, coax supports private independent networks that are not degraded by neighboring homes.

Medium	Shared	Channels	Physical Isolation Method	Without Physical Isolation
Power Line	Typically 6 to 8 homes at a transformer	1	Certified electrician installs expensive filters	Without filters, homes will interfere with each other
Wireless 2.4 ISM	Practically unlimited	1	Not practical	Interference is practically unlimited
Coax	2, 4 or 8 homes share a multitap	> 12	Simple \$1 passive filters between multitap and POE	Homes can coexist on separate frequency channels, without filters, preventing interference.

COAXIAL HOME PLANT CHARACTERIZATION

All communications on in-home coax today traverses between the POE and outlet, or from the headend to the room outlet (downstream) and from the room outlet back to the headend (upstream). There is no communication on in-home coax from room/outlet to room/outlet. In fact, the coax splitters used in homes are really directional couplers designed to isolate splitter outputs and prevent signals from flowing room/outlet

to room/outlet. This isolation is for two reasons:

- Reduce interference from other devices' local oscillator (LO) leakage
- Maximize power transfer from POE to outlets

With tuner LO leakages around -35 dBmV, and isolation output to output in splitters at around 20 dB, then the LO leakage will arrive at other devices around -55 dBmV, which is sufficient to have minimal or no degradation.

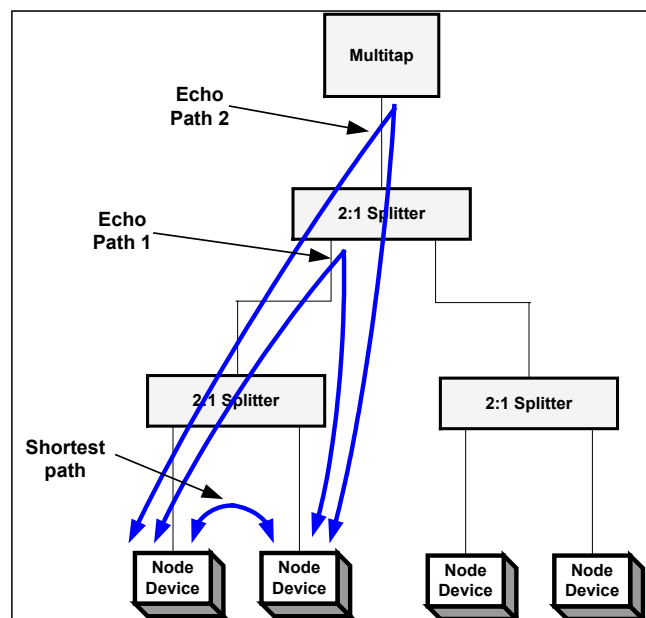
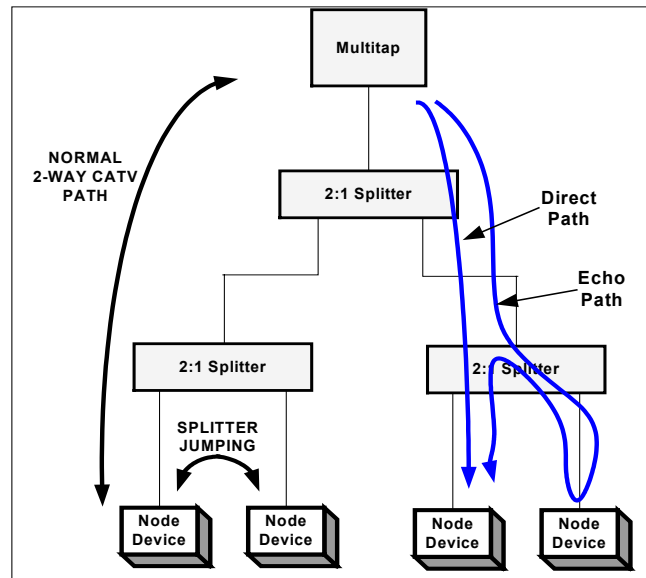
If true splitters were used instead of directional couplers, then the power loss through a 2-way splitter would be around 3dB higher in the normal direction.

Coax Echo Characterization

The following diagram shows the direction of signal flow in the normal and device to device directions on in-home coax. Note that device to device flow requires “splitter jumping” output to output.

Measurements of homes and splitters have shown that splitter output to output isolation is frequency selective and varies between 8 and 35 dB. Echoes in the home vary between around 12 dB to > 35 dB and so cover the same attenuation range, and often have close to the same attenuation, as the splitter output to output signal path.

The normal direction signal flow in a home creates echoes that are always delayed and attenuated. The following diagram shows that device to device signal flow creates echoes that can be more, less or equal in attenuation to the shortest path. Systems designed for POE to outlet echoes, such as DOCSIS and other single carrier modulations with QAM and a decision feedback equalizer, can not operate reliably in an environment from outlet to outlet where there are regularly zero dB echoes and seemingly non-causal echoes.

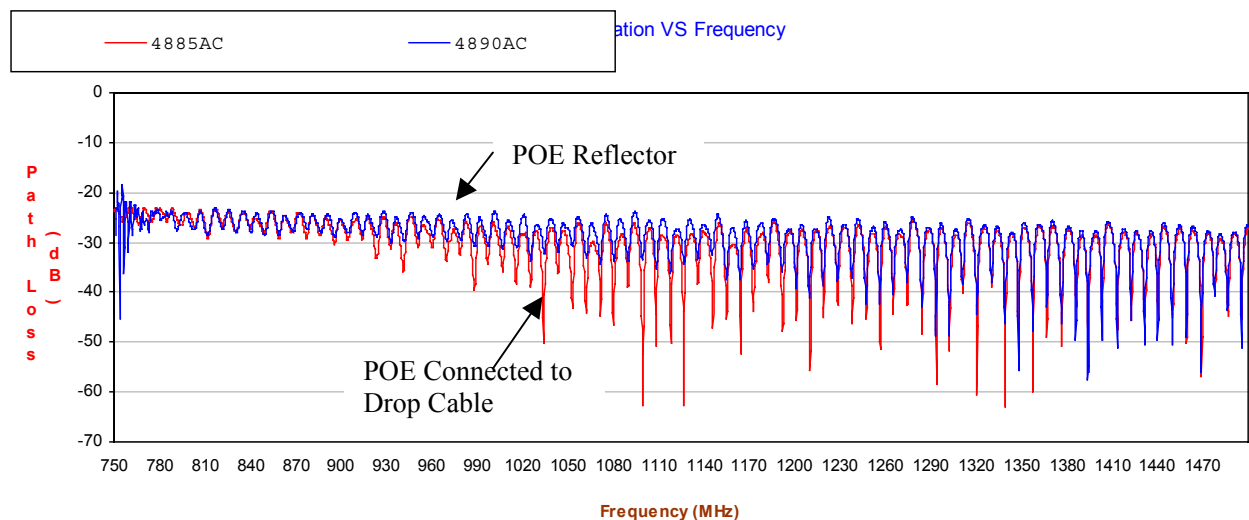
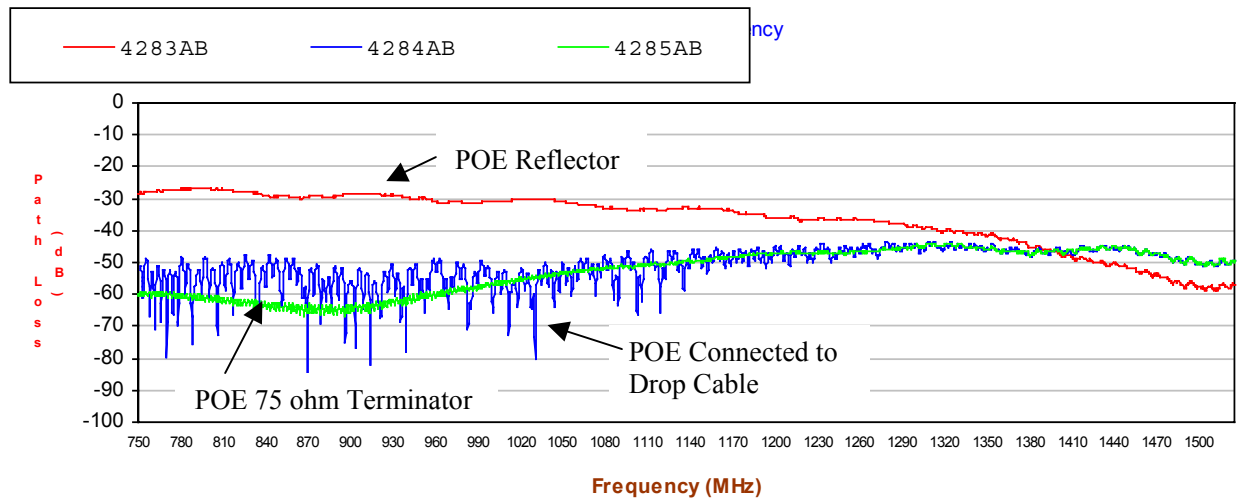


Two examples of characterized homes are shown in the following 2 figures. The first figure shows 3 scenarios: (i) POE connected to drop cable, (ii) POE terminated with 75 ohms, (iii) POE terminated with a reflector. Note that in the first figure the reflector case has less attenuation since the signal can go the normal path through the POE splitter and then perfectly reflect back instead of having to go output to output through the POE splitter and take the isolation attenuation. The 75 ohm

termination is very similar to when connected to the drop cable, but does not have the heavy notches from an echo off of the multitap.

From the first figure it might seem that installing a reflector at the POE removes the need to operate in zero dB and non-causal

echoes. However, the second figure shows that with or without a reflector at the POE, the home can have echoes generated from nested splitters, not the multitap, and thus still cause extreme echoes, which are manifested by the deep notches in the frequency response.



Coax Link Budget Characterization

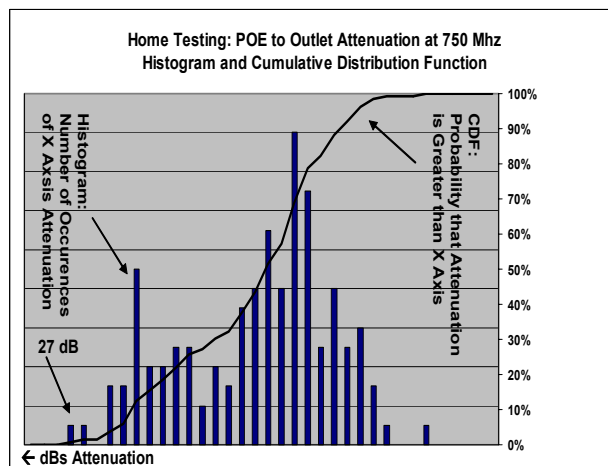
In addition to echoes that cause frequency selective attenuation, the overall average attenuation of an outlet to outlet link is substantial. Fortunately, the in-home cable

plant must have a minimum quality in order for existing services to work.

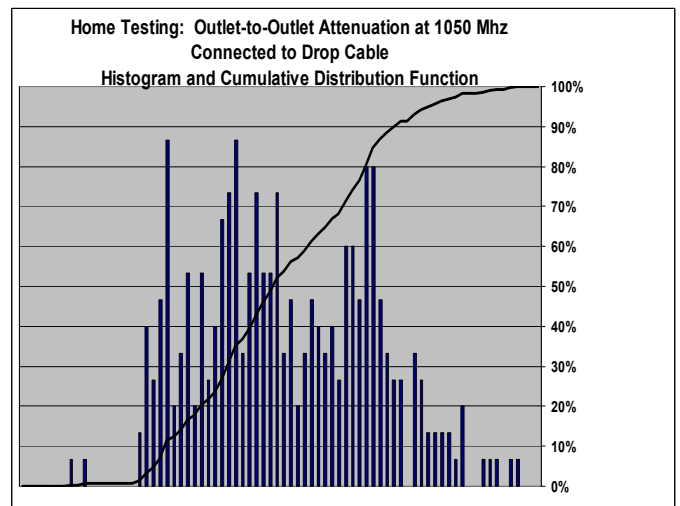
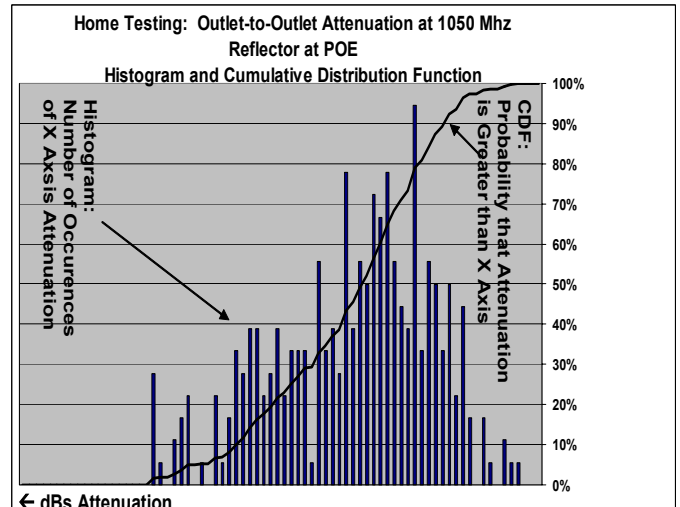
Existing digital and analog services tolerate a worst case attenuation of around 25 dB POE to outlet, at the highest frequency in the plant (750 Mhz, or 860 Mhz). Attenuation

beyond this results in non-operation of the target devices, such as STBs, and therefore the home network is not required to work either. This was corroborated with home testing and is shown in the following figure.

	Max Power at POE dBmV	Min Power at Device	Worst Case POE to Outlet Attenuation
Analog Video	+15dBmV	-10dBmV	25 dB
Digital Video or DOCSIS	+10dBmV	-15dBmV	25 dB



With a POE reflector the signal traverses outlet to POE, gets reflected and traverses POE to another outlet. The attenuation is thus around 50 dB at 750 Mhz or 860 Mhz. Without POE reflector the signal must go through the splitter isolation instead of twice through the normal attenuator direction, increasing the attenuation 1 to 28 dB. This was corroborated with home testing and the following two figures show outlet to outlet attenuation with and without a POE reflector at 1050 Mhz.



Home measurements showed that attenuation typically gracefully increases with high frequency out to around 1500 Mhz, after which it can become very attenuative and frequency selective.

Home Coax Signal Environment

The cable operator uses 5-42 Mhz for upstream and approximately 54-864 Mhz for downstream. The 42-54 Mhz band is for diplexer roll off in TVs, cable modems, and STBs and can not be effectively used by the home network without an unacceptable probability of interfering with such devices. Thus 5-864 Mhz can not be used and the

home networking must operate above 864 Mhz.

Coax Device Environment

The home coax plant has many existing devices in it including TVs, VCRs and STBs that are susceptible to interference from other signals on the coax. Even signals above 864 Mhz can degrade or overload existing devices so that there is a maximum power level that the home network signaling above 864 Mhz can be transmitted at. Characterization of over 130 existing devices showed that this non-interfering maximum power level is specifically frequency selective and varies in level from device to device. Double conversion tuners are in general less susceptible than single conversion tuners. A specific method of power control is required in order to meet the outlet to outlet link budget, yet not overload or degrade existing devices.

Home coax amplifiers are classified in two types: amplifiers in the drop side of the first splitter are termed “drop amps” and amplifiers in the home side of the first splitter are termed “inline amps.” Drop amps are outside of the home network signaling path and do not directly affect the home networking signal. Inline amps typically amplify in the direction of the devices from 50 to 864 Mhz and then roll off and can have significant attenuation above 1000 Mhz. Field tests and research indicate that around 2% of homes have inline amplifiers. Testing of inline amplifiers and homes indicated that in about half of the homes with inline amps the attenuation is not prohibitive and stays within the bounds of home attenuation described previously. This leaves half of the inline amp homes, or about 1% of homes, requiring remediation at the inline amplifier.

Remediation consists of either replacing the inline amp with one that supports signaling above 864 Mhz, or else putting a diplexer bypass around the existing inline amp. Such a bypass diplexer is low cost, and allows signals below 864 Mhz to operate as before, while signals above 864 Mhz are routed around the inline amp resulting in a passive network above 864 Mhz without high attenuation.

ENTROPIC’S c.LINK™ SOLUTION

Entropic has a production chipset and system solution for home networking triple play on coax. It includes a baseband controller doing MAC and Physical layers (BBIC), and an RF chip for conversion above 864 Mhz (RFIC).

The BBIC includes adaptive and modified multi-tone modulation over a 50 Mhz bandwidth, forward error correction, TDD burst generation and detection, mixed signal conversion, and an embedded processor that executes the TDMA MAC in software. Packet delivery across the link is guaranteed at 1e-9 BER without ever having a collision and thus not requiring retransmissions.

The RFIC includes LNA, PA, PLL for LO generation, quadrature up/down converter, and TDD controller to enable operation above 864 Mhz with up to 16 non-overlapping frequency channels.

Entropic’s overall system solution provides:

- Full mesh, peer to peer networking between all end points of a passive home coaxial cable plant
- Multiple, independent networks support via non-overlapping RF channels
- Physical layer data rate of 270 Mbps

- MAC layer (net) data rate of 135 Mbps
- 95% USA home MAC (net) data rate of >100 Mbps
- Co-existence with MSO cable services by operating above 864 Mhz
- Guaranteed/reserved bandwidth communications at 1e-9 BER
- Including isochronous traffic with jitter < 200 nsec
- Asynchronous communications at 1e-9 BER
- Including 802.1p 8 level prioritization
- Latency < 3 msec
- Link layer privacy encryption similar to DOCSIS privacy

Home Field Testing Entropic's Solution

Entropic has characterized and tested around 100 homes in southern California.

Entropic's testing showed that, for coax outlets that support existing digital and analog services, the 95% MAC (net) data rate is more than 100 Mbps from a physical layer rate of around 180 Mbps, and the 98 % MAC rate (excluding inline amplifiers) is more than 80 Mbps. This means that only 5% of outlets would not receive 100 Mbps or more MAC

rate, and 2% of outlets would not receive 80 Mbps or more MAC rate.

The field testing experience indicated that this puts the home networking reliability comparable to the probability that an outlet will even support existing digital and analog services. In other words, the home networking solution is consistent with MSO operational parameters.

CONCLUSIONS

A home usage model, system requirements and comparison of home mediums led to the conclusion that MSOs must run their whole home networked premium services video over either Cat-5 or coax. The unique coax home environment when communicating from outlet/room to outlet/room was then described and field characterization was presented. This environment was shown to include zero dBC and pre-causal echoes that make single carrier solutions such as DOCSIS and other QAM with equalizer solutions non-reliable. Entropic's solution and production chipset were described. Field testing and characterization in around 100 homes showed the Entropic solution meets a 98% reliability for the MSO, with reasonable remediation for the remaining few percent of outlets.

IP MULTICAST IN CABLE NETWORKS

Joe Godas¹, Brian Field, PhD², Alon Bernstein³, Sanjeev Desai³, Toerless Eckert³,
and Harsh Parandekar³

¹Cablevision Systems Corporation, ²Comcast Corporation,

³Cisco Systems, Inc.

Abstract

IP multicast is as an integral technology in networked applications throughout the world. Any network application involving the transmission of the same information to multiple recipients can benefit from the bandwidth efficiency of multicast technology. Multicast represents a key inflection point for the cable industry. While multicast is being used in cable networks today, two key new technologies—the wideband protocol for a Data Over Cable Service Interface Specification (DOCSIS) network and Single Source Multicast (SSM)—are expected to dramatically accelerate multicast deployment.

These technologies will help operators dramatically increase the operational efficiency of the Hybrid Fiber Coax (HFC) network, create a mechanism to accelerate the delivery of advanced services, leapfrog recent announcements of fiber-to-the-x (FTTx) deployments and service, and drive industry agendas for years to come.

This paper is jointly authored by Cablevision Systems Corp., Comcast Corp., and Cisco Systems, Inc. The paper describes multicast deployments at Cablevision and Comcast, highlights other multicast applications, and discusses key challenges. The paper proposes enhancements to DOCSIS specifications that should significantly increase multicast deployments in cable networks.

1. INTRODUCTION

IP multicast is a bandwidth-conserving technology that reduces traffic by simultaneously delivering a single stream of information to potentially thousands of subscribers. Multicast routing establishes a tree that connects a source with receivers. Multicast delivery sends data across this tree towards receivers. Data is not copied at the source, but rather, inside the network at distribution branch points. Only a single copy of data is sent over links that lead to multiple receivers, resulting in bandwidth gains. Multicast packets are replicated in the network at the point where paths diverge by routers enabled with Protocol Independent Multicast (PIM), and other supporting multicast protocols. Unlike broadcast, the traffic is only received and processed by devices that are listening for it.

IP multicast was developed in the early 1990s and was first deployed in education and research networks. About 1997, multicast was deployed on a large commercial scale when stock exchanges required a fast, efficient method to send market data to many subscribers simultaneously. For the past few years, multicast has gained wider acceptance as enterprises and service providers have realized the benefits of the technology.

Two multicast service models are deployed today:

- Any Source Multicast (ASM) is the original model introduced in 1990

(RFC1112) where an interested receiver of a multicast session notifies the network via Internet Group Management Protocol (IGMP) that it is interested in joining a specific group associated with that multicast session. The receiver then receives content sent by any source sending to this group. This model is targeted to support dynamic multi-source sessions like conferencing and financial trading. The standard protocol set in support of ASM is IGMPv2 or IGMPv3 for hosts to join a group and Protocol Independent Multicast-Sparse Mode (PIM-SM), together with Multicast Source Discovery Protocol (MSDP), for interdomain operations and rendezvous point (RP) redundancy. Support for IGMPv2 and ASM is covered in DOCSIS 1.1.

- Source Specific Multicast (SSM) is a more recent model in which an interested receiver of a multicast session specifies both the group and the source (or sources) from which it would like to receive content. The SSM model is superior for services where sources can be well-known in advance of the multicast sessions. The SSM model is achieved through the use of IGMPv3 which allows the host to specify both the group and the sources of interest, as well as the PIM-SSM which generates S, G joins in direct response to the IGMPv3 reports.

2. SAMPLE APPLICATIONS UTILIZING IP MULTICAST

This section highlights solutions in which IP multicast is an important element of the network. The section details deployments at Cablevision and Comcast.

Figure 1 on the next page depicts a sample multicast-enabled network.

This section also discusses multicast virtual private network (VPN) services. These services are offered primarily by telcos, but are of high value and interest to cable operators as well. The challenge for cable operators is to allocate enough spectrum and bandwidth to support these services. The wideband protocol for a DOCSIS network is the leading contender for providing this capacity. This section briefly describes the wideband technology and its relevance to cable operators as they converge IP services.

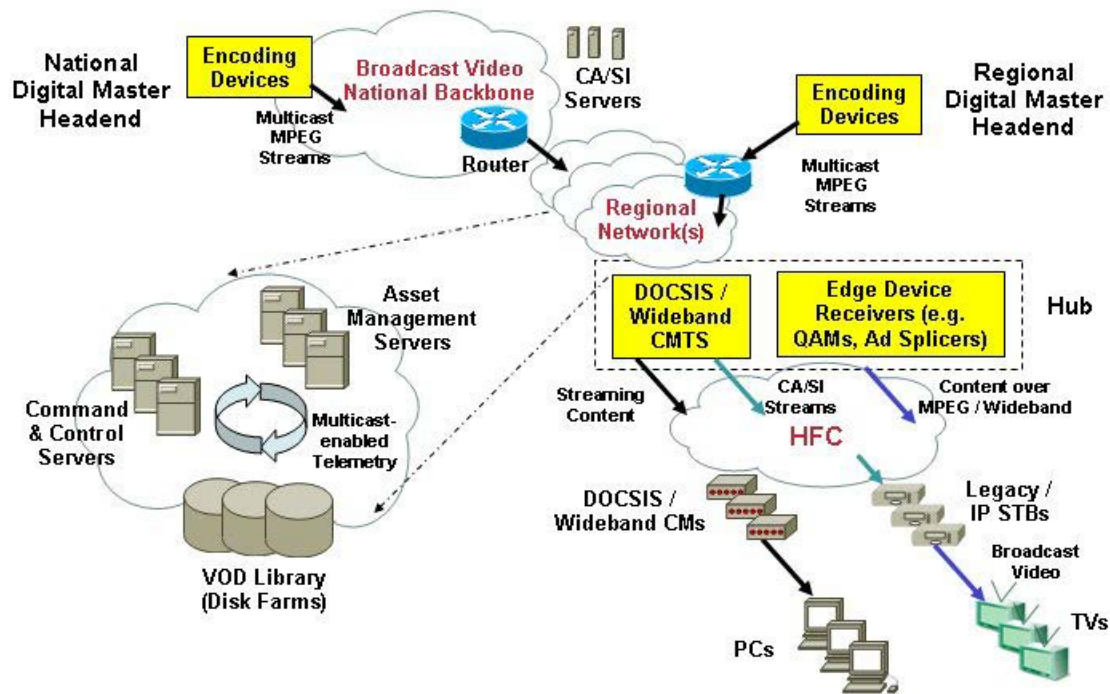
2.1. Digital Simulcast at Comcast

2.1.1 In Deployment

In today's broadcast video networks, proprietary transport systems are used to deliver entire channel line-ups to each hub site. These transport systems are often dedicated to broadcast, both digital and analog, video delivery and are not easily or economically extendable to other services.

By its very nature, broadcast video is a service well-suited to using IP multicast as a more efficient delivery mechanism. Comcast is in the process of moving its broadcast video service from a proprietary Baseband Video/Audio, IF, and DVB-ASI-based delivery system onto an IP network that is architected to support and deliver all Comcast-based services.

Figure 1 Network Diagram



The IP multicast delivery of broadcast video works as follows. Encoding devices in digital master headends, encode one or more video channels into a Moving Pictures Expert Group (MPEG) stream which is carried in the network via IP multicast. Devices at each hub site are configured by the operator to request the desired multicast content via IGMP joins. The network, using PIM-SM as its multicast routing protocol, routes the multicast stream from the digital master headend to edge device receivers located in the hub sites. These edge devices could be edge QAM devices which modulate the MPEG stream for an RF frequency or ad insertion devices which splice ads into the MPEG stream and then re-originate the ad zone-specific content to a new multicast group. Edge devices within the ad zone would use IGMP joins to request this ad zone-specific multicast content.

2.1.2 Futures

2.1.2.1 National Backbone

Comcast is in the process of deploying a backbone designed to support Comcast's specific service needs. This backbone will be multicast-enabled and be able to deliver broadcast video content to Comcast regional networks. Having a multicast-enabled IP backbone that is able to deliver broadcast video has a number of economic benefits, including the ability for the backbone to act as the backup origination location to the regional networks for core video channels. The cost of deploying high-quality video-encoding equipment in a backbone backup facility can be more easily justified as its expense is offset by reduced redundant encoding equipment needs in Comcast regional networks.

2.1.2.2 SSM

While multicast, as available today, is a useful technological solution for a number of cable service applications, there are areas in which further enhancements to multicast, multicast's interaction with the rest of network routing protocols, and with devices which participate in multicast, may be useful. One enhancement to multicast relates to using SSM instead of ASM. With today's

ASM /IGMPv2-based service with PIM-SM and IGMPv2, the complexity of IP multicast in the network is larger than necessary for applications with one or few (redundant) sources like DOCSIS Set-top Gateway (DSG) and Digital Simulcast. Migrating to PIM-SSM and IGMPv3 reduces this complexity, and thus, lowers the cost of operations. The challenge to adopting this technology lays primarily in edge device support like quadrature amplitude modulation (QAM) devices.

2.1.3 Challenges

Main challenges of this application are quality and availability. This translates into the applications and network redundancy design and failover times.

2.1.3.1 Better Overall Service Quality

As broadcast video is the core service offered by Comcast, the video service delivered via the IP transport network must be as good, if not better, than what is provided via existing legacy systems. Thus, it is critical that network and edge devices are highly available; the encoded video content is of very high quality; sufficient redundancy exists when a hardware failure occurs to enable the video service to recover quickly. Cable operators, therefore, have a number of design decisions to make regarding the level of device and network redundancy needed to

support the real-time broadcast video service requirements.

2.1.3.2 Redundant Sources

For critical channels or content, a cable operator may choose to have dual origination points in the network. If the primary facility becomes unavailable due to a catastrophic failure or unplanned maintenance, the content can be multicast from the backup location. The operator can opt to have this backup stream always "on" and immediately available to devices in the network. This mode results in fast service recovery, but at the expense of using more bandwidth in the network.

On the other hand, since losing the primary facility should be an uncommon event, upon losing the primary facility, the operator may opt to manually enable the backup feed. This reduces the amount of bandwidth needed in the network. The service recovery time, however, will be greater.

2.1.3.3 Fast IP Unicast Convergence

Since multicast relies on the underlying IP routing infrastructure to build the multicast distribution trees, the time to rebuild the multicast tree when a failure occurs in the network, is in part dependent on how quickly the unicast routing protocols re-converge. Only when the unicast routing protocols have converged can PIM begin to rebuild the multicast trees. Thus, for real-time multicast applications, it is important that the operator's network design enables—and corresponding unicast and routing architectures—support fast network re-convergence.

2.2 Multicast at Cablevision

2.2.1 In Deployment

2.2.1.1 System and Conditional Access Information Distribution to STBs

Cablevision currently uses IP multicast to drive the conditional access (CA) and system information (SI) carousels to their set top boxes (STBs). Their first advanced STB had a single out-of-band tuner which acted as both as an interactive and out-of-band (OOB) management interface. Cablevision originally supported DOCSIS and Digital Audio Visual Council (DAVIC) delivery mechanisms, but quickly adopted a DOCSIS-only approach once multicast robustness was demonstrated. Cablevision's newer STBs use DOCSIS to carry vital SI streams. Both CA and SI will be carried over into any DSG deployments that Cablevision evolves to in the future.

IP multicast is delivered via PIM-SM for all applications. The system and conditional access distribution itself does not rely on IGMP signaling from the STBs, but instead statically joins and forwards the traffic from the cable modem termination system (CMTS). Depending on the STB system, the packet flow ratios will be approximately 50 pps @ 300 kbits/sec or 110 pps @ 430kbits/sec.

In terms of separation, Cablevision's high speed data (HSD) customers are shielded from seeing these STB multicasts via standard DOCSIS cable modem filters which are established by the configuration from the cable modem's Trivial File Transfer Protocol (TFTP) boot file. These filters also prohibit unintended sources from hijacking or disrupting multicast flows. In addition, Cablevision prohibits IGMP from cable modems in the upstream direction through configuration on the CMTS. Only the groups of the streaming audio/video service

(described as follows) are allowed for subscriber cable modems of that service.

2.2.1.2 VoD Server Resource Management Telemetry

Cablevision's video on demand (VoD) resource management telemetry currently occurs over IP multicast, with a future eye on utilizing reliable multicast for content distribution to disk farms. The telemetry messaging is constant and averages close to 184 pps @ 458kbits/sec.

The command and control servers are a suite of servers that talk to clients and servers in the VoD cluster. The asset management component of the system must know what resources are available on each server in order to know if it has additional disk space and streaming capability. These status messages are carried over a proprietary multicast messaging system (developed by Seachange).

The basic network design consists of centrally located asset management and command and control servers that speak to remotely located VoD disk farms that communicate over the multicast network. Newer network design models and data transmission capabilities have enabled Cablevision to centralize the VoD disk farms. This in turn lends itself to a reduced need to route the VoD multicasts.

2.2.1.3 Streaming Audio/Video to Cable Internet Customers

This year, Cablevision is beginning a trial of real-time multicast video streaming to HSD users. This is meant to differentiate and add value to Cablevision's data service, promote loyalty, and reduce churn. Cablevision will start by porting selected Interactive Optimum (iO)—Cablevision's Video Service—content/functionality to Optimize Optimum Online (OOL)—Cablevision's cable modem

service—for use exclusively by subscribers to both iO and OOL services. Video and audio content will be offered. Rates for video will approach 500 kbps and 55 pps per stream. The same sparse-dense mode network used for STB's SI will be used. The encoding is Windows Media version 9 (WM9), but alternative encodings are also being investigated—specifically for future content over the wideband protocol for a DOCSIS network.

2.2.2 Futures

Looking to the future, Cablevision sees the use of multicast as a delivery mechanism for push-VoD models where content is streamed to a group of PCs for viewing at a later time. Cablevision's VoD libraries can be leveraged, in addition to third-party content providers.

Beyond that, Cablevision sees their switched broadcast architecture incorporating IP multicast to help drive the efficient delivery of popular content across the video backbone and down the respective QAM devices—be they traditional MPEG or IP over DOCSIS.

2.2.3 Challenges

While Cablevision has been successful with the systems and services it has deployed to date, the company needs to continue to refine its network strategy and fine tune its architectures to address ongoing changes and challenges. Some of these challenges include:

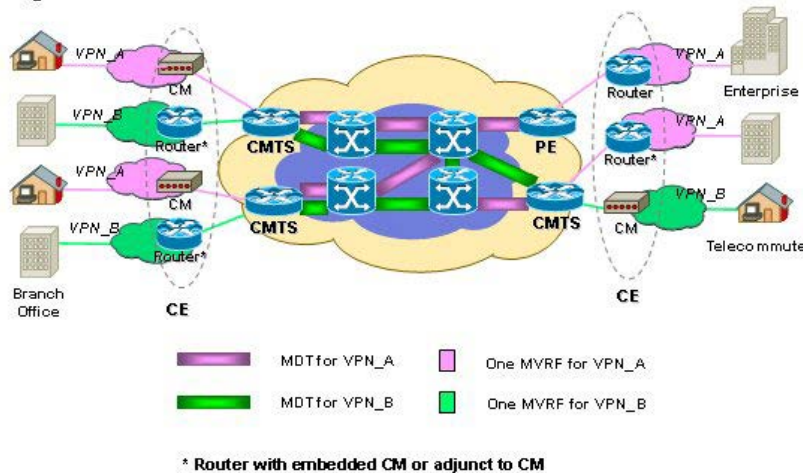
- DOCSIS 1.1 support is a necessity so that multicast flooding does not occur in a customer's home network.
- Enhancements to DOCSIS must be made such that multicast can be reliably scheduled and assigned a priority on DOCSIS segments.

- More multicast-aware customer premises equipment (CPE) gear: home routers and home wireless gear must continue to evolve to better support IGMP snooping, IGMP relay and firewall configurations that allow multicast streams to make it to intended destinations without allowing users to cannibalize their own experience. For example, a wired client on a home router should not be able to cause a multicast flood of his own wireless spectrum, if no wireless clients are requesting the flow.
- IGMPv3 and source-specific multicast (SSM) support
- More bandwidth: Cablevision must select their content carefully since there is a tight bandwidth budget with respect to the quality of the streams they want to offer. Cablevision is encouraged by the progress of the wideband protocol for a DOCSIS network and feels they will be able to exploit these opportunities further once larger backbones and modem contracts can be configured to handle cost-effective high-bandwidth services.

2.3 Multicast VPN Services

Commercial services over DOCSIS are steadily gaining traction in the cable environment—both in the U.S and abroad—because of strong revenue potential. One such service is VPN which allows businesses to connect multiple remote sites or devices over either a Layer 3 or Layer 2 VPN. Figure 2 depicts a multicast VPN service architecture.

Figure 2 Multicast VPN Services



For a Layer 3 VPN, the provider network is involved in the routing of traffic inside the VPN. A Layer 2 VPN provides a bridging transport mechanism for traffic between remote sites belonging to a customer. While these services are just gaining momentum in the cable world, they are quite pervasive in the telco world. In the telco VPN environment, enterprises have shown significant interest for native multicast support in the service provider's network. Current estimates are that ten to forty percent of VPN customers want IP multicast support in their VPN service to transport traffic for one or the other enterprise multicast application.

When VPN services are offered, multiple system operators (MSOs) will see the need to support IP multicast on these services. Typical enterprise multicast applications include NetMeeting, video conferencing, corporate communications, and finance-specific applications.

To support multicast over Layer 3 VPNs, each VPN receives a separate multicast domain with an associated multicast VPN routing and forwarding (mVRF) table maintained by the provider edge (PE) router. In the cable environment, the PE router can be a routing CMTS. The provider network builds a default multicast distribution tree (Default-

MDT) for each VPN between all the associated mVRF-enabled PE routers. This tree is used to distribute multicast traffic to all the PEs. For high-bandwidth multicast traffic that has sparsely distributed receivers in the VPN, a special MDT group called a Data-MDT can be formed to avoid unnecessary flooding to dormant PE routers.

IP multicast can also be supported in Layer 2 VPNs via IGMP and/or PIM snooping in the provider's network. L2VPN services can be provided by configuring the CMTSs for point-to-point tunneling or for multipoint bridging. Depending on the configuration, snooping takes place on the external Layer 2 aggregation device or on the CMTS. Based on the snooped messages, the multicast traffic can be forwarded only to those customer edge (CE) devices that are interested in that traffic, versus flooding it to all the CE devices.

Security and data privacy are of primary concern in a VPN environment. The service provider network, including the CMTS, must be able to distinguish between multicast sessions that belong to different VPNs. On the shared cable downstream, packets belonging to separate VPNs must be encrypted using separate BPI keys. Since group addresses used within different VPNs can overlap, multicast support in VPNs can be complex without the right support in DOCSIS.

2.4 High-Speed IP over Cable with the Wideband Protocol for a DOCSIS Network

Cable operators are now entering the third phase of service convergence as they increasingly add IP video services to existing data and voice IP offerings. Service delivery requirements are rapidly evolving. A significant percentage of traffic will shift from broadcast video to per-user streams as deployment of network VoD services enables consumers to move to a user-controlled “watch whenever” viewing paradigm.

In the short term, the transition to per-user video streams is largely taking place in the MPEG domain. VoD and personal video recorder (PVR) services are being delivered to conventional STBs via MPEG transport streams. Over the longer-term, more of the video content will be delivered via an end-to-end IP infrastructure—directly to televisions and PCs in the home. Therefore, the infrastructure deployed must be capable of evolving into an all IP network. The wideband protocol for a DOCSIS network, along with multicast and IP Version 6 (IPv6), are key ingredients of this evolution.

The wideband protocol, which is under evaluation for inclusion in pending DOCSIS 3.0 specifications, allows cable operators to make the leap to IP video faster and cheaper than telco companies. The technology supports bonding multiple channels to allow cable operators to add downstream channels, independently of upstream channels. The technology enables operators to leverage previously deployed DOCSIS CMTSs and take advantage of declining prices for external edge QAM devices. It will allow operators to use the same edge QAM pool for both data and video services. The technology provides plenty of bandwidth for multiple standard-definition digital and HDTV channels, IP telephony and data offerings, with a capacity of up to 640 Mbps.

3. TECHNICAL CHALLENGES

While there are numerous challenges overall, this section concentrates on what we consider to be the top two challenges:

- Issues with ASM
- Limitations in DOCSIS 1.1

3.1 Issues with ASM

Three basic issues with ASM and its protocols exist:

3.1.1 Address Assignment

In ASM, only one application can use a group address G at a time. As long as multicast applications only need to run between participants within a single administrative entity, this is manageable. A single administrative entity can construct an address plan of RFC1918 type IP multicast group addresses (from 239.0.0.0/8). But this requires operational coordination which adds cost. If on the other hand, multicast traffic is to be transported across domains—for example from a content provider onto one or more cable operator or telco network—then coordination of IP multicast group addresses becomes an almost unsolvable problem.

3.1.2 Denial of Service Attacks

The ASM service model is prone to attacks by unwanted sources, because receivers do not specify which source(s) they want to receive traffic from. While it is possible in a walled garden network to provide additional network-based access control, the operational cost of such control rises as more and more multicast applications are deployed in the network.

3.1.3 Complexity of Provisioning and Operations

Unlike older multicast protocols, the PIM-SM/MSDP protocol provides efficient delivery of traffic and high availability. This comes at the cost of adding many protocol elements which increase the complexity of the network.

Amongst these are:

- Placement of RPs, operations, and troubleshooting of RPs
- Operations of BSR or Auto RP protocols for RP redundancy
- Alternatively, static configuration of RPs and set up of MSDP-mesh groups for anycast-RP
- Operations of MSDP between administrative domains
- Troubleshooting of PIM-SM protocol elements such as RPT/SPT switchover and register tunnel encapsulation

3.2 Limitations in DOCSIS

Current DOCSIS specifications define several hooks for enabling multicast on the RF. These include:

- Baseline Privacy Interface (BPI) extensions that allow encryption of multicast sessions
- IGMP snooping in the cable modem (CM) that is used to trigger the BPI exchange for multicast.

The purpose of IGMP snooping is to restrain multicast traffic and specify how a host can register a router to receive specific multicast traffic. These specifications leave a wide range of issues unaddressed. These are discussed below:

- Aliasing of traffic: According to RFC1112, aliasing of traffic may

happen because only the lower order 23 bits of an IP multicast address are mapped to a multicast Ethernet address. For example, a CM configured to receive traffic for group 224.1.2.3 will accept traffic for 239.1.2.3. This is particularly an issue with VPN and SSM support.

- Limited support for multicast protocols: DOCSIS 1.1 does not have IGMP support for IGMPv3 and SSM, Generic Attribute Registration Protocol (GARP), GARP Multicast Registration Protocol (GMRP) or Multicast Listener Discovery (MLD).
- IPv6: No support for IPv6 multicast since IPv6 has a dedicated control plane for multicast.
- PacketCable Multimedia (PCMM): PCMM does not yet define how multicast is to be supported.
- Lack of explicit tracking of multicast listeners: Because of IGMP v1/v2 report suppression, the CMTS cannot track which hosts are actually listening to a given session and cannot support fast-leave for multicast sessions.
- Quality of Service: QoS for Multicast flows is not defined
- Routed networks on the CPE side: The network cannot easily support routers connected in the CPE network that run PIM instead of IGMP.
- JOIN acknowledgment: There is no way for a CM or CPE to make sure a multicast session is successfully activated since there is no explicit acknowledgment in IGMP V2.

4. SOLUTIONS

4.1 SSM

4.1.1 Solving ASM Issues

SSM solves several problems with the ASM service model:

- Network-wide group address allocation: In SSM, the multicast group G does not need to be unique over the network because only (S,G) channels need to be unique. Groups can be reused.
- DOS attacks: Receivers will only receive traffic from the source which was explicitly indicated in their IGMPv3 joins.
- Simplified operations: In PIM-SM, a receiver host joins to a group G. The network builds a delivery tree towards an RP (the RP tree) and sources register to the RP via an encapsulation tunnel. Then, the RP joins to the source to receive traffic from the source and sends it down the RP tree. Once the router connected to the receiver sees packets from a new source S arriving on this tree, it joins to this source via the Shortest Path Tree (SPT)—also called the (S,G) tree

In contrast, with PIM-SSM, the IGMPv3 (S,G) report from the host allows the router connected to the receiver to bypass all the initial steps involving an RP and start out immediately by establishing the SPT for (S,G).

SSM with IGMPv3 and PIM-SSM is an evolutionary technology because PIM-SSM is a subset of PIM-SM. Routers that support PIM-SM also support PIM-SSM. Applications supporting SSM with IGMPv3 will also work in an existing PIM-SM IGMPv2 network, because IGMPv3 is

automatically backwards compatible with IGMPv2.

4.1.2 Challenges of SSM Deployment

The challenges in deploying SSM are adoption and support of IGMPv3 with (S, G) receiver reports in applications and appliances; for example, STBs, PCs or QAM devices.

SSM mapping can be used as a transition strategy. In SSM mapping, the router connected to receivers is seeded with the source address belonging to groups G. While the receivers only send IGMPv2 reports for groups G, the router itself adds the source address and then continues to use PIM-SSM.

4.2 DOCSIS 3.0 Multicast Proposal

Multicasting on the RF can save bandwidth on the RF interface. Currently, however, DOCSIS RFI specifications do not fully address multicast. A DOCSIS 3.0 specification proposal has been submitted to CableLabs to address current DOCSIS limitations on multicast.

It suggests the following framework:

- Multicast flows are signaled in the same way that unicast flows are—through registration or DSx message exchange. They use the same TLV unicast flows use—an admission control function to keep track of the fact that multicast flows do not consume additional bandwidth once the first one is established.
- The multicast control plane handling is moved to the CMTS.

The list below outlines current issues with DOCSIS multicast support, and explains how the multicast proposal addresses these issues:

- Aliasing of traffic: Current DOCSIS specifications support only RFC1112 mapping of multicast addresses. With this mapping, two separate groups can be mapped to the same MAC address. The proposal recommends setting a multicast media access control (MAC) address from a CMTS-allocated pool of multicast MAC addresses outside of the RFC1112 MAC address range. The CM can later replace this locally assigned address to a standard RFC1112.
- Limited support for multicast protocols: Current DOCSIS specifications use "IGMP snooping" to detect that an IGMP was sent from the CPE. By moving the IGMP control plane processing to the CMTS, the system is not limited to "snooping" multicast and inherent problems associated with snooping. Instead, the end point that was supposed to receive the multicast—the CMTS—is the one responding to it.
- PCMM: Currently there is no PCMM definition on how multicast can be handled. Since the multicast proposal to CableLabs treats multicast as unicast in terms of flow definition and setup, then PCMM will tie seamlessly into this framework.
- Limited monitoring on the CMTS: An explicit signaling for multicast flow set up will allow for deterministic tracking of multicast users, instead of relying on "report suppression"
- Quality of service (QoS) definitions: Current DOCSIS specifications have a rich set of methods to define QoS. However, these are tied to a specific modem. The proposal allows these definitions to be re-used for multicast as well.
- Routed networks on the CPE side: if PIM is running between the CMTS and a customer's router, the CMTS can

trigger a multicast DSx based on the PIM state machine.

- JOIN acknowledgment: Currently, there is no explicit response to an IGMP v2 JOIN. If the JOIN triggers a DSx message exchange, the DSx-RSP will return specific error codes if the multicast session cannot be established.

5. SUMMARY AND CONCLUSIONS

IP multicast has a wide range of applications for current and future cable operations.

Cable-specific applications (to QAM devices or STBs) include:

- Digital simulcast of live TV via IP multicast (e.g., Comcast)
- Switched video/TV broadcast (dynamic overprovisioning to save bandwidth)
- DSG/proprietary STB system information and crypto key distribution (e.g., Cablevision)
- VoD server resource management (e.g., Cablevision)

"Enterprise" applications include:

- Reliable content/software distribution / preprovisioning with Pragmatic General Multicast (PGM) or other multicast transport to VoD or Web servers
- VoIP: multicast music on hold, voice conferencing (Hoot & Holler)
- Enterprise corporate communications, video conferencing, corporate event broadcasting, and training
- Financial applications including stock trading and market data distribution
- Retail including warehouse-distributed applications (e.g., with TIBCO middleware) (typical drivers for VPN customers asking for multicast)

DOCSIS 1.1 applications include:

- Live audio/video streaming to HSD customers (e.g., Cablevision)
- L2/L3 VPN services: delivering "enterprise" multicast applications

Wideband Protocol for DOCSIS Network applications include:

- Higher bandwidth applications, more customers/content, and HDTV
- Key to migrate cable-specific applications to IP (with an IP STB that supports the wideband protocol)

While cable operators may start with as little as one application, they will likely need to support multiple applications over time. This leads to the conclusion that IP multicast will be one of the core capabilities of a cable operator's IP network for the foreseeable future and will help operators unleash the full power of their HFC network and architecture.

But before this promise can be fulfilled, there are a number of items that must be considered and decisions to be made. The two most important network technologies that cable operators must consider in conjunction with IP multicast are SSM and the wideband protocol for a DOCSIS network.

SSM can be deployed today. The challenge is to ensure it is supported in applications and appliances such as STBs.

The wideband protocol for a DOCSIS network will expand cable operator service profiles in the IP/data arena. Its challenge, in conjunction with IP multicast, is to ensure improved support for several elements in IP multicast (like SSM), as outlined in our DOCSIS 3.0 proposal.

6. REFERENCES

- [ASM] "Host Extensions for IP Multicasting", RFC1112, S.E. Deering. Aug. 1989
- [IGMPv3] "Internet Group Management Protocol, Version 3", RFC3376, B. Cain, S. Deering, et. al., Oct. 2002
- [PIM-SM] "Protocol Independent Multicast - Sparse Mode (PIM-SM) Protocol Specification (Revised)", draft-ietf-pim-sm-v2-new-11.txt, PIM-WG, Oct. 2004. Notes: 1. The PIM-SM RFC2362 is obsolete. 2. PIM-SSM is a subset covered by this specification, too.
- [SSM] H.Holbrook, B.Cain, "Source-Specific Multicast for IP", draft-ietf-ssm-arch-*.txt, Sept. 2004
- [IGMPv3SSM] "Using IGMPv3 and MLDv2 for Source-Specific Multicast", draft-holbrook-idmr-igmpv3-ssm-08.txt, H. Holbrook et. al, Oct. 1, 2004
- [DOCSIS 3.0 Proposal] "Multicast Proposal for DOCSIS 3.0", Submitted to CableLabs by Cisco Systems, Inc., Dec. 10, 2004
- [DOCSIS 3.0 Proposal] "Wideband Proposal for DOCSIS 3.0", Submitted to CableLabs by Cisco Systems, Inc., Dec. 10, 2004

IP VIDEO TRANSPORT SOLUTIONS FOR CABLE OPERATORS

Mark Davis¹ and David Brown²

¹S. V. Vasudevan, ²BigBand Networks, Inc.

Abstract

With content and service expansion a constant in the cable industry, operators have had to adapt their plants for increasing capacity demands, while also containing costs for maximum return on investment. This has driven several cycles of evolving transport techniques to deliver content from centralized sources to distributed subscriber bases.

IP technologies have extended new promise and seen initial implementations that achieve superior economics and functionality for data and voice, and increasingly video content. The openness and flexibility of IP enables rapid advances and cable operators can now contemplate the economic utilization of fiber distribution at both regional and national levels. There are several courses to consider for more widespread utilization of IP technologies for the distribution of core cable programming, and this paper considers their relative advantages and challenges, while recommending commitment to a highly flexible infrastructure able to adapt to emerging opportunities.

VIDEO TRUNKING EVOLUTION

Analog satellite was the first “trunking” technology used to cost effectively deliver multiple video programs over long distances from a single point to multiple destinations of nationwide headends. But evolutionary drivers proceeded quickly, as cable operators continued to see capacity

requirements explode along with operator consolidation in the 1980s and 1990s.

Needs emerged to retool the metropolitan network architecture from standalone headends to headend-hub architectures. Most cable operators maintained wireless practices by leveraging CARS band microwave for video trunking between these facilities, until fiber optic transport became the new industry standard capable of transmitting analog video over relatively long distances without sacrificing picture quality. This optical transport method utilized FM and AM modulation techniques.

The AM super-trunk eventually emerged as the low-cost reliable solution and is still used for many hubs today. The advantage has been that signals could be processed or modulated at the central headend and then distributed to hubs up to 30 miles away over a single mode fiber. 1550nm optical transmission technology evolved with higher output lasers and EDFA amplifiers which pushed distances up to 90 miles. A simple low-cost optical to electrical receiver kept space requirements to a minimum and allowed retransmission to optical nodes in neighborhoods, commonly known as hybrid fiber/coax today.

However, system consolidation and the pressure to cut more operating costs and add more channels continued through the 1990s, which in turn drove a new digital transport technology. Systems emerged using proprietary encoders integrated with SONET-like transport, allowing delivery without the picture quality challenges or distance limitations of AM super-trunks. In

addition, small decoders converted digital video streams back to IF frequencies and allowed the use of small low-cost IF-to-RF upconverters. Limitations included only 16 channels per fiber capacity and proprietary techniques that constrained interoperability and concentrated market power with particular vendors. SONET and even ATM vendors did make several attempts to solve these issues and were able to capture a small number of sites, but primarily failed due to the costly and heavy floor space requirements of baseband to RF conversion or modulation required at all hubs.

Thousands of SONET-oriented terminals were deployed at distant digital hubs. The industry then launched digital programming to the home in the late 1990s, requiring expensive digital video headend and video processing equipment. AM super-trunking remained a compelling solution through the transition to MPEG video because of its cost effectiveness. However, limitations on transport distances and functionality such as program grooming and hub-based local insertion still required some use of transport terminals until now.

IP TRANSPORT EMERGES

Recent widespread deployments of VOD (Video On-Demand) has driven new requirements for transporting large quantities of video streams and high volume QAM modulation at the hub or edge of the network. This has prompted reconsideration of the industry's video transport techniques. Early VOD deployments relied on costly ASI transport and distributed VOD servers. Technology vendors responded with new low-cost, high-density QAM modulators with integrated IP-to-MPEG-2 decoders that used Gigabit Ethernet as the transport layer. These edge QAMs were deployed using a variety of optical connections including

DWDM (Dense Wavelength Division Multiplexing) and CWDM (Course Wavelength Division Multiplexing) technologies. New standards-based SFP (Small Form-Factor Pluggable) optical transceivers have allowed transport and edge QAM vendors to drastically cut the costs associated with both CWDM and now DWDM transport. In addition, GigE over next generation SONET as well as IP routing solutions costs also plummeted. These events opened up a host of new standards based, low-cost, and flexible video transport solutions.

Pushing digital programming deeper into the network has improved picture quality and reliability, lowered operating costs by consolidating headends, and provided capacity to meet ever growing programming requirements. Catalyzed by VOD requirements, utilizing IP technologies for this deep digital transport has opened new opportunities for the industry.

As digital broadcast channel capacity continues to increase, an operator has to consider the economics of transitioning away from traditional 1550nm super-trunks or continuing to expand proprietary digital transport. With the drastic decreases in costs outlined above, transitioning all digital programs to MPEG-2 over IP/GigE can be achieved at similar costs to proprietary digital system expansion, with significant benefits from these open and highly functional standards.

A more recent driver of digital transport evolution and bandwidth management requirements is the introduction of digital simulcast. This is the digital encoding of all analog channels, along with QAM modulation of the traditional analog tier that is then simultaneously, along with a decoded version of the analog tier, transmitted to

subscribers. Encoding of all remaining satellite delivered analog programs along with hundreds of PEG (Public access, Educational and Government) channels is required for digital simulcast.

Simulcast allows operators to start the migration to an all-digital network and provides multiple near-term benefits such as use of less expensive, all-digital STBs (Set-Top Boxes), and improved picture and audio quality especially to high-end televisions. Simulcast also drives the transition of analog commercial ad insertion to an all-DPI (Digital Program Insertion) solution. DPI servers can be centralized and upgraded to provide GigE outputs, allowing DPI or digital splicing anywhere in the network, including hub-based ad zones. Simulcasting leverages the benefits of Gigabit Ethernet and IP transport techniques, advancing their predominant establishment in cable networks for all media and services.

Industry consolidation is continuing to drive the need to further consolidate and convert headends to hubs in larger metro areas and even fiber-connected rural areas. Many smaller headends in sparsely populated areas are now connected with leased or owned fiber for high speed data and telephony applications.

As the number of subscribers served by a headend increases, operators must consider adding a second redundant headend on the backbone in order to ensure the highest reliability or network availability, while delivering the lowest possible cost service. MPEG-2 over standards-based GigE and IP allows operators to utilize MPEG aware switching platforms to perform automatic failover or improve fault tolerance at the program level.

As transport trends progress, the metro concept is expanding and the benefits of nationwide terrestrial transport based on GigE, IP and optical technologies is becoming apparent. Several operators now have substantial wide-area fiber resources, or can cost-justify lambda leases from long haul fiber providers and consolidate all video, voice and data traffic onto a common backbone. Smaller operators are also faced with all the same challenges and have started forming partnerships to build a common channel IP video backbone fed by redundant primary headends.

FUNDAMENTALS OF MPEG-2 AND IP VIDEO NETWORKING

MPEG-2 transport as standardized by ISO 13818-1 has established itself as the de-facto protocol for the carriage of broadcast digital television services. Even the newer advanced video codecs such as H.264 can be carried over MPEG-2 transport. A coded video or audio frame is fragmented into several MPEG-2 transport packets, which are a fixed 188 bytes in length. A 13-bit PID (Packet Identifier), present in every MPEG-2 transport packet, identifies the elementary video and audio packet streams that comprise a broadcast television program. By using different PIDs to separately distinguish different elementary streams, an entire program can be multiplexed and carried as an SPTS (Single Program Transport Stream). Expanding this methodology, several programs, each with their own unique video and audio PIDs, can be broadcast as a unified multiplex, known as a MPTS (Multi-Program Transport Stream).

To carry MPEG-2 transport over IP networks, a further encapsulation step is required to place the MPEG-2 transport packets in an IP networking envelope. UDP

(User Datagram Protocol) is most commonly used to carry both broadcast and narrowcast MPEG-2 traffic. This has advantages for real-time content like video over widely known TCP (Transmission Control Protocol) by eliminating the need for packet-receipt acknowledgements back to the transport source. Since the 24 byte UDP header represents additional transmission overhead, it is prudent to carry as many MPEG-2 transport packets in a single UDP message as possible. Typically,

seven MPEG-2 packets are placed in a UDP message. This number is chosen because it represents the maximum UDP message size that fits into the maximum 1,500 byte payload size that is dictated by the 802.3 Ethernet framing format. While larger UDP message sizes are possible, limiting the overall payload to less than 1,500 bytes guarantees that the message will not be subjected to any unnecessary IP fragmentation procedures that could be asserted by lower-level protocol stacks.

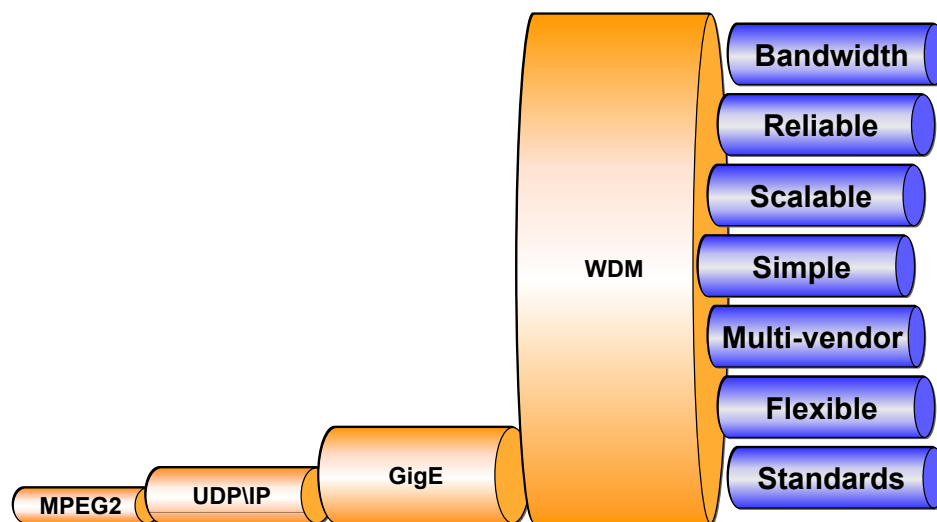


Figure 1. Multiple digital encapsulation techniques can be openly nested within each other providing a range of attractive benefits for transport of content and services.

MPEG-2 is sensitive to jitter, which is variation in delay in the arrival of transport packets. This is because the transport stream carries timing information that is used by the receiver (e.g. an STB) to faithfully decode and regenerate the baseband program of interest. In an end-to-end digital television system, a 27 MHz system clock is typically locked to the incoming baseband video stream, although it is also possible to originate a stream with a free-running clock. This 27 MHz system clock drives a counter that generates a 42-bit PCR (Program Clock Reference). PCRs are inserted into the transport stream at a regular rate (at least 10

times/second), and are in essence real-time snapshots of the counter in the encoder. Decoders receive and extract PCRs for the particular program they are decoding and use the values and appropriate filtering to drive a VCXO (Voltage-Controlled Crystal Oscillator) on the receive side to accurately regenerate the 27Mhz clock, which is then used to regenerate the baseband video timing. This transmission model allows a decoder in a consumer's home to regenerate a video clock that is locked in phase to the encoder that originally compressed the program, which could be in a satellite uplink facility thousands of miles away.

Minimizing PCR jitter is important to maintaining a robust and high-quality end-to-end signal. ISO 13818-1 specifies a maximum PCR jitter of 500 ns, but this value specifically does not include jitter that can be caused by UDP/IP encapsulation and network transport by IP or other protocol. It is not uncommon to observe UDP/IP-encapsulated network traffic with tens of milliseconds of PCR jitter, almost two orders of magnitude above that maximum limit specified by MPEG.

To cope with this high network jitter, most video-aware networking components de-jitter the stream by intelligently correcting the PCR value in the transport packet and/or shaping the flow of the packet traffic as it transits through the device. Without an effective de-jittering mechanism, excessive PCR jitter can cause packet deliveries that violate the buffer models specified by MPEG, and more importantly can prevent the decoder from accurately regenerating the 27MHz clock required to reconstruct the baseband signal. Consumer televisions that encounter this type of impairment will generally be unable to lock to the corrupted “color burst” waveform that precedes the delivery of a field of NTSC video – the most common symptom of this situation is that the video picture will be unable to render any chrominance information and becomes “black-and-white”.

Another approach to de-jittering is the use of RTP (Real-time Transport Protocol), as specified by IETF RFCs 1889 and 2250. RTP is a popular protocol for streaming media applications, and has mechanisms to support multi-source applications such as videoconferencing. A UDP message would be the data payload of an RTP packet. An RTP header contains a timestamp that can be used to recover and restore packet timing

between a network sender and receiver. However, jitter that is induced by the UDP encapsulation of MPEG-2 transport packets cannot be recovered with RTP timestamps, and RTP adds another 12 bytes of overhead to the overall message size.

NEW VIDEO BACKBONE REQUIREMENTS

IP capabilities are becoming increasingly applicable for video traffic in cable networks at an opportune juncture of both service expansion opportunities and competitive threats to the industry. Customer tolerance to network outages will quickly diminish now that the competitive landscape is changing with telcos and satellite video providers vying for market share. Given this risk, any video network architecture design must have very high availability, fault detection and rapid recovery features.

Increasing channel counts along with more HD programming will continue to drive the need for additional capacity. Channel expansion must be easy to implement while minimizing costs.

Pressure to minimize headcount expenses while driving high network availability, will require networks to remain simple to provision, operate and easy to troubleshoot. Minimizing the number of appliances and network elements such as small-profile but limited-purpose “pizza boxes” minimizes the amount of training and device dependencies. Complex routed networks can also be difficult to troubleshoot when a number of multi-service traffic related challenges like denial of service attacks appear on the video network. Safeguards and sound QoS (Quality of Service) practices must be implemented to avoid these challenges.

Now that the billion dollar access network upgrade is complete, there is financial pressure on MSOs to drive free cash flow by limiting capital spending. This will drive efficient use of existing infrastructure like fiber and existing data transport elements. WDM technologies are an attractive way to extract maximum returns on existing specialized network element investments while allowing growth and graceful transition to an all IP video transport network. There is no need to do a fork-lift upgrade of all transport devices with WDM. Upgrades can happen when service specific requirements dictate as with the case of video transport expansion requirements.

The renowned flexibility of IP expands cable operator choice in how transport networks are configured and operated. A combination of business and technology drivers are positioning the industry to consider much more widespread transport implementations than the classic headend-hub metro connection with nationwide digital distribution. Several implementations of this are available for MSO consideration. Robust and flexible infrastructure proves key to maintaining best-of-breed options over time while implementing the best techniques to address current needs.

PASSIVE WDM NETWORKS

Passive WDM systems are now readily available for a variety of network transport systems. Passive WDM utilizes industry standard ITU (International Telecommunication Union) grid lasers that are tuned to a specific wavelength. CWDM typically uses bands S, C and L normally with 20nm channel spacing and center wavelengths at 1491nm, 1511nm, 1531nm, 1551nm, 1571nm, 1591nm, 1611nm. CWDM wavelengths can start as low as

1270nm. The same techniques are used for DWDM networks as well. However, the wavelength spacing is much tighter at .4nm (50Ghz), .8nm (100Ghz) and 1.6nm (200Ghz) from 1525 to 1615nm which drives up component cost for both lasers and passives.

Most optical transceivers are now designed to comply with the SFP industry-standard MSA (Multi Source Agreement). Ethernet switches, routers and SONET multiplexers have all now incorporated SFP standards. Low-cost SFP repeaters can provide 3R regeneration (re-amplify, reshaping and retiming) to overcome loss and dispersion limitations of long links.



Figure 2. SFP modules dramatically improve price-performance and flexibility of optical networking.

SFP transceivers feed optical combiners and splitters called optical mux-demux passives. These low-cost devices separate wavelengths or lambdas while avoiding high insertion losses. These devices use GRIN (GRAdient INdex) lenses or thin filters that are normally packaged together to filter two or more wavelengths. GRIN lenses focus light through a precisely controlled radial variation of the lens material's index of refraction from the optical axis to the edge of the lens. By gradually varying the index of refraction within the lens material, light rays can be smoothly and continually redirected towards a point of focus. This

allows a GRIN lens with flat or angle polished surfaces to collimate light emitted from an optical fiber or to focus an incident beam into an optical fiber. End faces can be manufactured with an anti-reflection coating to avoid unwanted back reflections

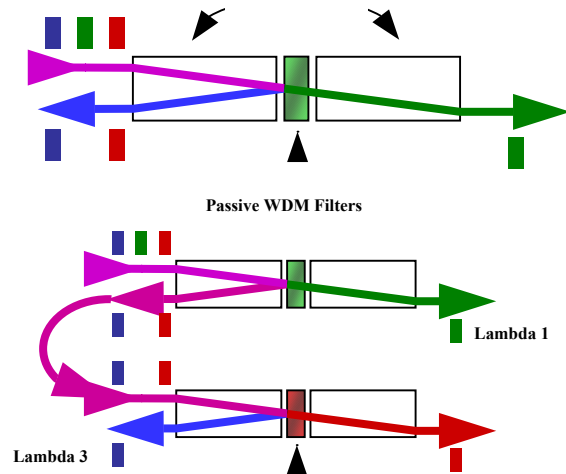


Figure 3. GRIN assures stable and reliable optical networking performance through precise lens characteristics.

Most metro networks have larger fiber counts and need fewer channels to transport a full broadcast video lineup to remote hubs. A four lambda system carrying four Gbps of bandwidth over 100 km costs under \$2,500 per GigE. Two GigE links can carry over 480 standard definition video programs. Low-cost regeneration at drop sites or mid span repeaters can be deployed to overcome long distances.

The real-world, large metro-based CWDM network represented in figure 4 provides digital video transport between two redundant headends and 24 digital hubs serving over 1.3 million homes. Two GigE links on a single fiber provide capacity for up 48 38.8Mbps MPTS broadcast video multiplexes. Each hub has dual route diverse links fed from two separate headends which prevents fiber related outages.

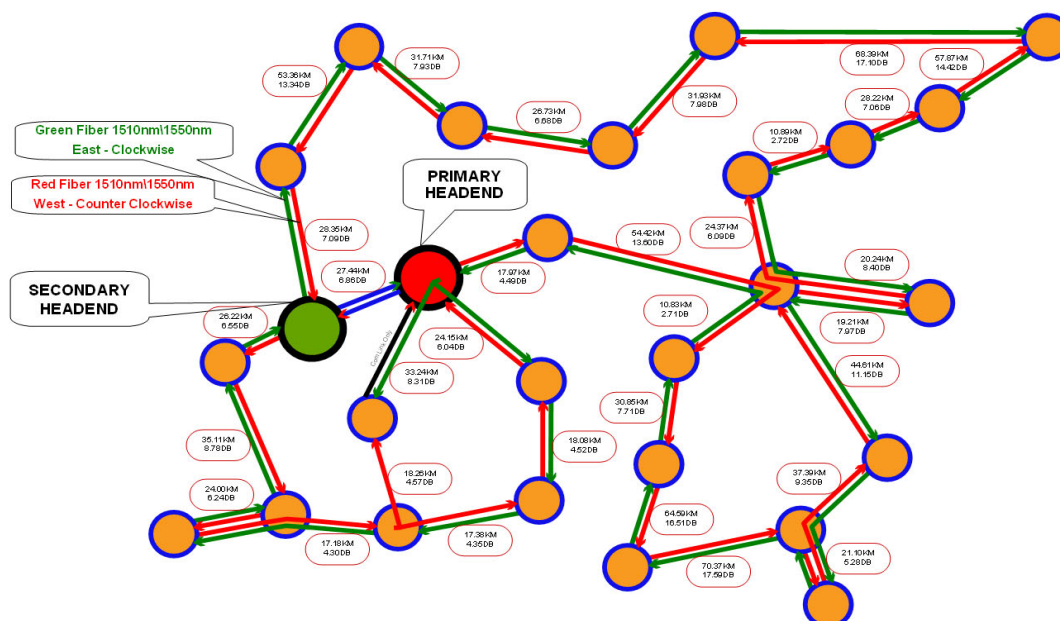


Figure 4. A passive CWDM network utilizing primary and secondary headend feeds and a bidirectional transport architecture for service assurance.

The biggest advantage of passive CWDM is simply the lower cost. Wide channel spacing enables less stringent laser performance requirements and lower cost optical passives. Small, integrated SFP optics require less rack space and power consumption. Passive DWDM systems are best when more than eight Gbps are required and fiber resources are limited. Passive WDM networks are very simple to provision and expand when required. In short, there are less active or moving parts to break which results in better availability. WDM systems are able to transmit multiple bit-rates and protocols, which allows the use of optimized application devices. SONET for CBR (Constant Bit Rate) T1 and DS3 traffic, Ethernet switches and routers for IP traffic and MPEG-2 routers for GigE over IP traffic allows systems to deliver more specialized features for a longer useful life.

There are a few limitations when using CWDM including the lack of wideband amplifiers. SFP CWDM repeaters work well when there are a limited number of wavelengths to retransmit but can be costly when dozens of lambdas need amplification. DWDM allows the amplification of multiple wavelengths with a single device.

ACTIVE WDM NETWORKS

Active DWDM systems work similar to passive networks but have finely tuned lasers and cross-connect features. These devices are known as transponders and muxponders.

One challenge when operating passive DWDM networks is crosstalk and balancing of optical power levels. Active DWDM networks use precise transponders for managing power levels and 1+1 optical redundant switching. High performance lasers also allow transmission over ultra-

long fiber spans up to 200 km. Active transponders can also be equipped with electrical multiplexers for aggregating multiple GigE feeds into a single 10 Gbps lambda. Active DWDM networks are best when fiber resources are extremely limited and the maximum amount of bandwidth per lambda is required.

The downsides of active DWDM include added costs and more active or moving parts to fail.

GIGE OVER SONET

SONET networks have been the workhorse of local and inter exchange carriers for over a decade. SONET was optimized for high availability or lifeline CBR transport such as DS3 and DS1 traffic. Several layers of hardware protection ensured five nines of reliability. Today's next-generation SONET network elements now have integrated lower cost layer 2 Ethernet functionality and GigE interfaces. These interfaces can be used in a one-way cross-connect mode that can efficiently transport MPEG-2 MPTS over IP to distant hubs. Figure 5 shows how one-way drop and continue cross-connects for unidirectional video traffic operate while still providing a protection path in the case of a catastrophic fiber cut.

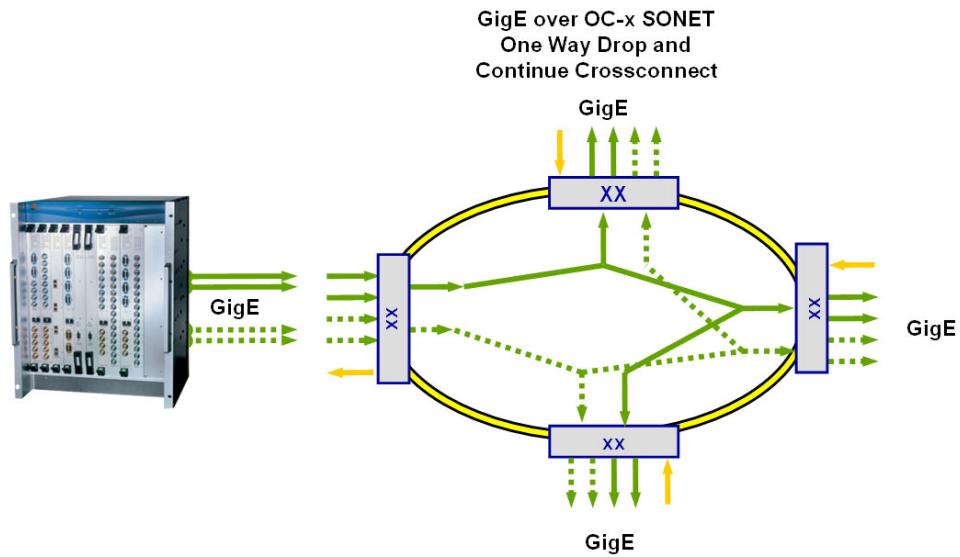


Figure 5. IP/GigE and SONET techniques can be combined for reliable and economical transport while leveraging prior investments and practices.

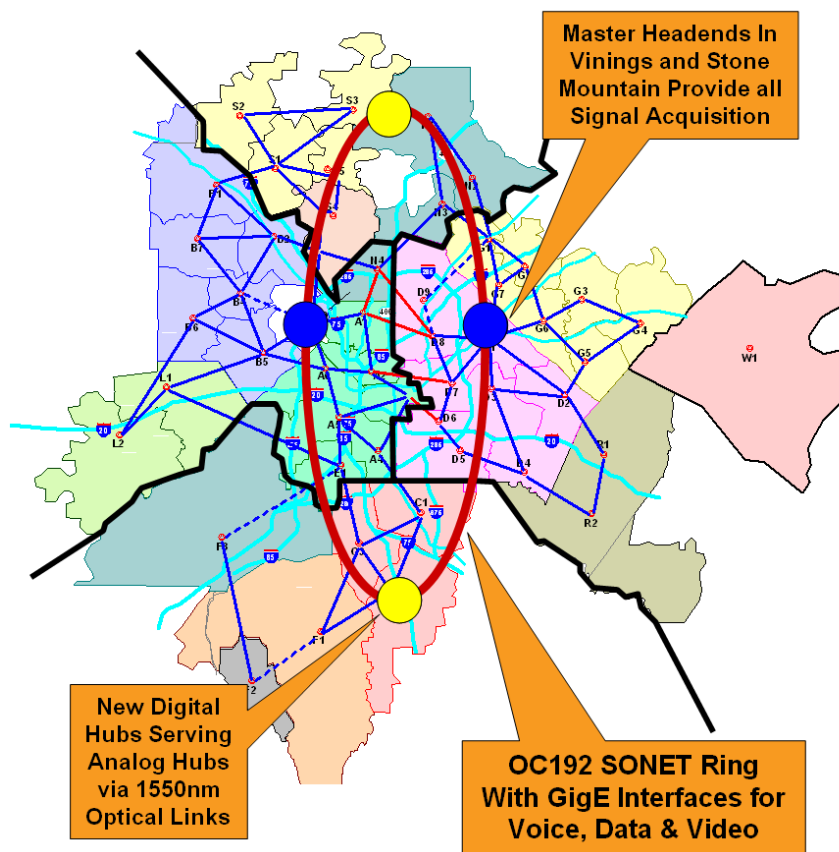


Figure 6. real-world network leveraging multi-use SONET rings.

OC-192 SONET terminals have had significant drops in price and size over the last five years. Many operators already have IP and voice traffic traversing these rings. Excess SONET capacity can easily be

The biggest challenge associated with SONET networks is scalability to meet the growing increases in pure IP-based traffic. Many new IP-based switching and routing platforms are catching up to the high performance track record of SONET-based solutions at far better price points. A recent SONET approach called RPR (Resilient Packet Ring) is now emerging. If prices can continue on a downward path, RPR may challenge classic IP routing with several of the advantages of packet-based routing.

CONVERGED REGIONAL AREA NETWORKS

One way to design the network is using a routed IP backbone. Cable modem and VoIP (Voice over IP) services have been running over a routed IP network for years, while video has traditionally been sent over a parallel network. There are a number of reasons that MSOs have not been sending video traffic over an IP routed network until now:

- Until recently, IP routers did not have sufficiently robust QoS mechanisms to ensure lossless, low-jitter video delivery.
- Video equipment vendors only recently added support for IP-based video transport in their products.
- Video uses a tremendous amount of network bandwidth, and the Gigabit Ethernet capacity has been traditionally insufficient to benefit from the convergence of multiple services.

provisioned for a robust MPEG-2 over IP/GigE broadcast video transport link. SONET terminals by far have the best track record for network availability performance.

Today, however, these issues have been addressed. All major router vendors include QoS mechanisms, IP transport is rapidly becoming a standard transport option on video equipment, and 10 Gigabit Ethernet is available at very competitive price points. As a result, MSOs can now make use of the same IP routed network for all of their video, voice and data services.

IP networks can use unicast (one-to-one) or multicast (one-to-many) transmission. While VoIP and data services currently running on MSO IP backbones is unicast, video services are multicast in nature. As a result, video services carried over IP networks are typically multicast. It is critical to have a high-performance multicast-enabled IP network in order to carry video services.

A key benefit of running video services over IP routed networks is the fact that routing protocols (e.g., OSPF and BGP-4 for unicast; PIM-SM and MBGP for multicast) make the network highly dynamic and robust. Rather than having operators statically define paths from a video feed in a headend to a hub site, the IP network will automatically calculate and transmit along the least-cost routed path. If there is an equipment failure, the routers will dynamically discover the failure and re-route around it – without manual intervention.

Today's multicast IP networks generally use ASM (Any Source Multicast). In ASM networks, receivers (e.g., a splicer/groomer in a hub site) use a protocol called IGMP (Internet Group Management Protocol)

version 2 to join a multicast group. It is called “any source” multicast because the IGMPv2 protocol does not specify the sender, only the group. After sending the IGMP request, the network re-configures itself to make that video feed available to the requesting device.

When using ASM for multicast IP, a multicast routing protocol is generally required to ensure that the multicast traffic is routed to all the correct network destinations; PIM-SM (defined in RFC 2362) is typically used. In PIM-SM, all traffic initially is transmitted via a router designated as the RP (Rendezvous Point). Because all multicast traffic is sent via the RP, an appropriate multicast stream can be found for each multicast group. The disadvantage of this approach is that the RP

can be a bottleneck and potentially a single point of failure.

A new approach called SSM (Source Specific Multicast, RFC 3569) eliminates these challenges. With SSM, receivers must specify both the multicast group *and* the IP address of the multicast source. IGMPv2 does not support this capability, so IGMPv3 – a new version of the IGMP protocol – is required in order to support SSM functionality. A key advantage of SSM is that because the receiver specifies a source in advance, a rendezvous point is no longer needed. This eliminates a potential bottleneck and point of failure from the network. The short term challenge with SSM is that most existing devices do not yet support IGMPv3.

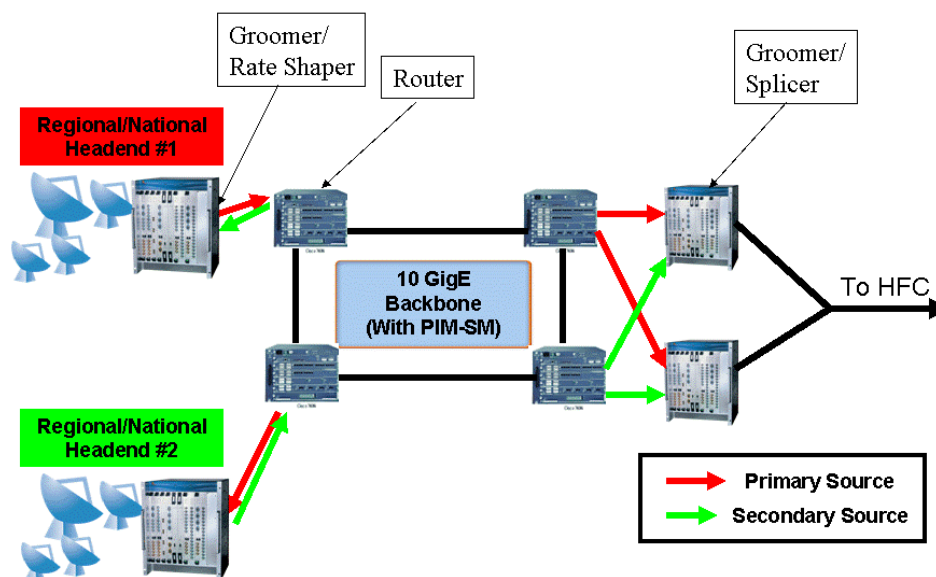


Figure 7. Incorporation of IP routers is a substantial modification of cable architectures, allowing high flexibility to support various functionalities and future directions.

An advantage of moving to an IP routed network is that the MSO can converge video and other services such as data and VoIP services on the same IP backbone. With only one network to manage, operations costs are reduced – and bandwidth can be shared among the various applications.

Running a converged network requires QoS mechanisms to be in place. Video traffic is very sensitive to jitter and packet loss, while VoIP is sensitive to latency and packet loss. In IP networks, the DiffServ protocol (defined in RFC 2474) is the most commonly used QoS mechanism. A DSCP (DiffServ CodePoint) is a six-bit field in the IP header of each packet – specifying the queueing behavior for that packet. Video and VoIP packets will be marked as high priority, and will be transmitted ahead of data packets. Using DiffServ, there can be lossless, low-latency, low-jitter transmission of video and VoIP even in the face of data congestion. Another protocol called MPLS (Multi Protocol Label Switching) can be used in conjunction with DiffServ, providing additional control over which traffic flows over which links. The disadvantage of MPLS is that it can add significant complexity to network operations.

Detractors debate the merit of converging video and other services on a single IP backbone, citing a number of concerns:

- On a converged network, video services can be brought down by denial of service attacks.
- Video programming is always-on, and of fixed bandwidth, reducing the statistical gain made possible by the sharing the network with other services.
- A converged network increases the complexity of network configuration, requiring careful QoS and traffic engineering to ensure that jitter and

packet loss do not negatively impact video quality. This may increase the operating expenses significantly enough to negate any benefits of a converged network.

Several mitigating practices can address these challenges.

A primary network design consideration is security. This should include location of encryption which can be central, at the acquisition site, or at the edge, in the RAN (Regional Area Network). From a practical standpoint, each RAN must have its own DNCS (Digital Network Control System), or something similar, for handling communications with STBs in the region, and they already have CA (Conditional Access) systems in place. As such, the content need not be encrypted centrally.

However, particularly in a converged network, content owners – fearful of internet hackers somehow getting free, unencrypted programming – may require that all content be encrypted. As such, IPSec or a similar protocol may be used to encrypt traffic between the central acquisition site and regional or metropolitan headend. This remains a challenge, since IPSec devices generally do not have sufficient performance to encrypt such a large quantity of data.

Because a large number of headends (potentially every system owned by a given MSO) rely on feeds from the central acquisition headend, guaranteed resiliency is a must. To maintain such resiliency, several actions should be implemented. Programming identical to that at the primary headend must be available from a secondary headend that is in a geographically separated location. In a converged network, video traffic must be given bandwidth guarantees to ensure the network is fully non-blocking

for video traffic. Video quality monitoring capabilities must be in place at each regional or metro headend. Program-level redundancy should be used, and quality monitored continuously on a per-program basis.

Ensuring video quality across a national network is not easy. In general, today's cable networks may have an end-to-end latency of 50 ms. As MSOs move to national networks, latency will increase, additional jitter will be introduced, there will be multiple paths from source to destination. In addition, there are more points at which congestion and packet loss may occur.

For MPEG traffic, the maximum acceptable jitter is 500 ns. Encapsulating MPEG in IP alone introduces significant jitter, since up to seven MPEG packets are encapsulated in one 1,500 byte IP (over Ethernet) packet, producing queueing delays that introduce far more jitter than the specification allows. As a result, it is necessary to de-jitter the traffic at the receiving side. Poor de-jittering may result in stutter in the picture, and lower video quality.

To ensure high quality, the video portion of the network must also be virtually free of packet loss.

CONCLUSION

Passive WDM, active WDM, GigE-SONET and converged RAN are all worthwhile considerations for cable operators who want to leverage IP technologies for transport of core programming services. This has further advantages in to varying degrees enabling convergence with IP techniques used for emerging services like VOD as well as voice and data offerings. As a result, management capabilities and economics are improved. Benefits can be extended to national fiber distribution that revolutionize program sourcing practices.

However, all of these techniques do bring their relative advantages and challenges. Selecting for short-term optimization can position an operator less well as needs change. Key to this is building current IP infrastructures on robust, flexible and scalable platforms that are programmable for ongoing implementation of the best and most recent transport techniques.

Mark Davis
Vice President,
Network Solutions
mark.davis@bigbandnet.com

S. V. Vasudevan
Chief Architect and
VP Systems Engineering
vasu@bigbandnet.com

David Brown
Director,
Product Marketing
david.brown@bigbandnet.com

BigBand Networks, Inc.
475 Broadway
Redwood City, CA 94063
650 995 5000
<http://www.bigbandnet.com>

MAKING BUSINESS SENSE OUT OF THE WIDEBAND PROTOCOL FOR A DOCSIS NETWORK

John T. Chapman
Cisco Systems, Inc.

Abstract

To date, cable operators have enjoyed an upper hand in the competition to deliver high-speed services because of their network capacity and bundling strategy. In the U.S., telcos' attempts to counter this advantage by reliance on digital subscriber line (DSL) service have been somewhat successful in attracting new subs. But because of data throughput restrictions and the carrier's reluctance to enter the video space, these services did not threaten cable operators' advantage. However, with recent announcements of fiber-to-the-X (FTTX) and metro Ethernet architectures in the U.S., and deployments of advanced DSL, fiber networks and true triple play services elsewhere in the world, the playing field looks much more level and the true battle is ready to begin. The situation is even more intense in the Asia Pacific, Japan, and European regions where competitive broadband services are showing a growing traction with residential customers. So how can cable operators respond to these threats and the telco promise of 25, 50 or even 100 Mbps and greater broadband service to the home?

Fortunately for the cable industry, the answer does not lie in replicating its \$80+ billion investment to add new physical capacity on top of existing networks or matching the \$10 to 20 billion telcos will invest in fiber-based IP services. The answer lies in unleashing the full power of the cable industry's existing hybrid fiber coax (HFC) networks.

With sixty percent of today's HFC spectrum being used to carry less than two percent of its data capacity, the goal is to take the available HFC bandwidth and data capacity and use them more efficiently. By simply modifying the connectivity of the backbone to the plant, operators will achieve ten times, 100 times, or even 1000 times of today's cable data capacity and at significantly lower price points than existing per-port broadband costs.

The technology that makes this possible is the wideband protocol for a Data Over Cable Service Interface Specification (DOCSIS[®]) network. The technology promises to leapfrog the telco fiber strategy, dramatically alter the communications industry competitive landscape, and unlock more upside revenue potential for the cable industry than the original specifications.

TECHNOLOGY OVERVIEW

In order to fully understand the potential of existing HFC networks, let's first look at an example of how HFC plant capacity is used today. The chart below details the current utilization of a cable network supporting 100k HHP.

HFC Plant Capacity

Case Study

100K HHP HFC Plant
500 HHP per
Fiber Node

Service Group (SG) is
the no. of fiber nodes
with the same
service offering

BW Used Today
÷

Total Available BW
=

Efficiency

Service	# Channels		Digital BW	SG Size	Capacity	
Analog	79	60%	x 3.75 Mbps	x 1 SG =	0.3 Gbps	2%
Digital	43	33%	x 38 Mbps	x 1 SG =	1.6 Gbps	8%
VoD & DOCSIS	9	7%	x 38 Mbps	x 50 SG =	17.1 Gbps	90%
Total	131 Channels				19 Gbps	
Capacity	131 Channels		x 38 Mbps	x 200 SG =	1000 Gbps	

1.9%

An analog video channel occupies (in this example) 6 MHz of RF bandwidth, but it only carries one video signal. Also, because of the broadcast nature of the network, the 6 MHz of bandwidth effectively serves a single service group: the entire 100k HHP. The bottom line is that sixty percent of the HFC network's bandwidth currently yields less than two percent of the network data capacity. The spectrum utilization efficiency increases in the case of digital video, where typically up to 10 or 12 video channels can be transmitted in a 6 MHz RF channel, but cable operators are only starting to scratch the surface of their investment potential with adoption of video on demand (VoD) and DOCSIS.

The take away from this exercise is shocking. With 1000 Gbps available in today's cable networks, less than two percent of its total capacity is used today!

Let's look at the way current DOCSIS technology makes use of HFC capacity. In North America, DOCSIS 1.0 and 1.1—collectively known as DOCSIS 1.x—support 30.34 Mbps to 42.88 Mbps (approximately 27 Mbps to 38 Mbps usable throughput) in a single downstream RF channel, and 320 kbps to 10.24 Mbps (approximately 300 kbps to 9 Mbps usable throughput) in a single upstream RF channel. DOCSIS 2.0 moved the bar higher yet, allowing symmetrical data transmission by increasing the upstream raw data rate to as much as 30.72 Mbps in a single RF channel. DOCSIS 2.0's downstream technology though has remained the same as DOCSIS 1.0 and 1.1.

With the adoption of emerging Internet applications such as music and video download or interactive on-line gaming, per-channel data throughput is rapidly becoming a bottleneck. In order for cable operators to meet subscribers' growing bandwidth requirements, the downstream and upstream data rate limits will need to be increased. A

number of approaches have been suggested to achieve this goal. The accompanying table summarizes the pros and cons of these methods.

Data Throughput Enhancement	Pros	Cons
Use higher order modulation (e.g., change from 64- to 256-QAM, or 256-QAM to 1024-QAM)	Low silicon risk Minimum MAC change	Requires very clean plant Higher carrier-to-noise ratio necessary for same BER
Different PHY layer technology (OFDM, Wavelet, etc.)	Improved per-channel data capacity	Incompatible with existing cable modems Requires fork-lift upgrade Not proven technology in cable networks Extensive MAC changes
Increase cable network's operating RF bandwidth (e.g., from 50-860 MHz to 50-1000 MHz)	Low risk technology Increase available downstream spectrum and channel capacity	Requires major rebuild or upgrade Capital expense to upgrade plant
Change from subsplit to midsplit band plan (e.g., reverse spectrum 5-108 MHz rather than 5-42 MHz)	Increase upstream spectrum RF bandwidth Low-risk technology	Requires changing diplex filters in all actives Requires changing actives if diplex filters are hard-wired Must balance/sweep align all actives after diplex filter mods Need to make

Data Throughput Enhancement	Pros	Cons
		sure reverse amplifier modules/circuits work to 108 MHz; if not, replacement required Loss of some downstream RF spectrum
Node splits	Backward compatible Well-understood Being done by many cable operators now Increases effective RF bandwidth per subscriber	Capital expense to upgrade plant (materials and labor)

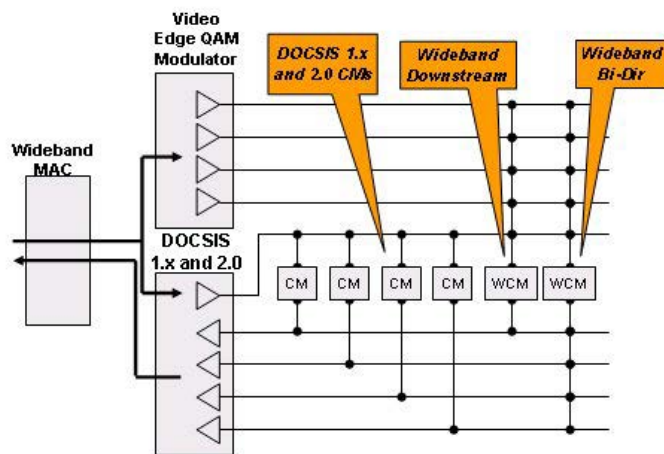
This table was adapted from “Next Gen Full-Service, All-Digital HFC Network Beyond DOCSIS 2.0”, a paper presented by John Eng at the Society of Cable Telecommunications Engineers’ 2004 Conference on Emerging Technologies. All of the alternatives in this table have a substantial impact on CapEx and/or OpEx. Of those listed, splitting nodes is the only one being done to any extent by the cable industry.

Given the spectrum inefficiencies of one analog TV channel per 6 MHz of RF bandwidth, along with the maximum per-channel throughput limitation in DOCSIS 1.x and 2.0 technology, it’s clear that cable network capacity has room to grow. That growth doesn’t require increasing the available RF spectrum, but rather making better use of the spectrum.

One solution, originated by the author in 2001 and under development by Cisco Systems® Inc., is known as the wideband protocol for a DOCSIS network. It solves bandwidth and throughput problems by logically bonding multiple RF channels together to form a wideband “channel,” and works equally well in the downstream or upstream.

Using the wideband protocol, data is striped across multiple quadrature amplitude modulation (QAM) channels, yielding a single logical channel—a wideband channel—the aggregate capacity of the individual QAM channels. The number of QAM channels logically bonded is dynamically configurable, providing the flexibility to increase the aggregate channel capacity with simple software configuration. For instance, by bonding four 256-QAM channels, one would obtain a wideband channel with a data rate of 171.52 Mbps (~152 Mbps). If downstream data were striped across 24 256-QAM channels, the result would be a wideband channel with a data rate of 1.029 Gbps (~912 Mbps). This approach allows operators to overcome the current DOCSIS per-channel downstream limit, without changes at the physical (PHY) layer. The same 64-QAM or 256-QAM modulation formats used today by DOCSIS 1.x/2.0 can be used for each of the channels in the wideband bundle) and without touching at all the network topology.

The following figure shows a high-level view of wideband technology overlaid on a DOCSIS 1.x and 2.0 network.



The wideband protocol is designed to be backwards-compatible with existing DOCSIS 1.x and 2.0 networks. It also delivers some of the benefits of CableLabs' Modular CMTS (M-CMTS™) architecture, such as separation of media access control (MAC) and PHY, or the use of edge-QAM devices using today's technology.

As we will see in the next section, it not only addresses the increasing data capacity and throughput demands, but it does so leveraging current cable modem termination system (CMTS) technology and existing edge-QAM devices to provide a lower cost per port than current DOCSIS QAM technology. In addition, we will see how cost efficiencies are possible because additional downstreams may be added independent of upstreams and existing RF spectrum can be used without impacting cable plant infrastructure costs.

ECONOMIC BENEFITS OF THE WIDEBAND PROTOCOL FOR A DOCSIS NETWORK

This section outlines the economic and business benefits of the wideband protocol for a DOCSIS network. Benefits can be effectively categorized as follows:

Increase Revenue

- Grow subscriber base: Attract new customers, retain existing subscribers, and fend off competitive threats with higher speed services
- Increase average revenue per user (ARPU): ability to offer higher tier services and expand service portfolio to allow operators to benefit from higher average revenue per user

Optimize CapEx Investment

- Fully exploit today's HFC plant potential: no network upgrades are required to take advantage of the throughput increase offered by the wideband protocol
- Leverage existing CMTS platforms: implementations will enable existing CMTS platforms to support the new protocol, lowering incremental capital expense
- Reduce downstream port cost: leverage lower prices of QAM technology and enable operators to use existing edge QAM devices
- Eliminate port under-utilization (stranded ports): add upstream and downstream ports independent of one another to accommodate traffic requirements
- Improve network efficiency: increase the subs/port ratio by taking advantage of the enhanced statistical multiplexing characteristics provided by a larger pipe
- Scale to future requirements: wideband components are designed to meet rapidly changing subscriber demands which in turn will protect initial investment.

Minimize Operational Costs

- Backwards-compatible with DOCSIS 1.x/2.0: fully leverage existing DOCSIS 1.x/2.0 provisioning systems and operational processes to simultaneously serve wideband customers

Deploying the wideband protocol for a DOCSIS network represents an evolution of the existing cable infrastructure, while offering a revolution in service capabilities.

This section examines in more detail the previously listed benefit categories.

Increase Revenue:

Grow Subscriber Base

It's a fact that the speed of Internet connectivity is at the top of the purchasing decision criteria for broadband subscribers. Today, in many metropolitan areas around the world, consumers are already being offered the following broadband choices:

Access Technology	Provider-	Offered Throughput
Competitive Carrier - Fiber Overbuild		100 Mbps +
Incumbent Carrier - DSL		25 Mbps
Cable Operator		8 Mbps
Wireless/Satellite Provider		2 Mbps

In Japan, YahooBB! has found tremendous market success with its 100 Mbps up to 1 Gbps residential service, demonstrating just how beneficial it is to be the highest-speed provider.

Within the U.S., cable operators have enjoyed success in the broadband battle. Thus far, roughly two thirds of broadband customers use cable modems. Competitors are

aggressively attacking this position by investing in a combination of copper-based, fiber-based, and wireless technologies.

These facts point to the need for cable operators to increase their service offerings to meet real world capacity requirements and competitive positioning. Operators must show they can continue to meet customer needs going forward. While DOCSIS 1.x and 2.0 have accomplished this to a certain point, network data throughput demands will continue to eclipse this capacity. Wideband technology and the bonding of multiple channels give operators the ability to far exceed today's data offering and put cable networks in a decidedly favorable position. This added capacity gives operators tremendous flexibility in defining new service models and revenue opportunities. Operators can provide increased choice and the ultimate on-line experience. They can expand beyond established subscriber demographics to attract a broader set of users.

The wideband protocol for a DOCSIS network offers the ability to deliver the highest service throughput (100 Mbps and higher) and unmatched service selection. This puts cable operators in a leading position in the highly competitive broadband market.

Increase Average Revenue Per User (ARPU)

Several analyses have shown the positive effect of a tiered service offering versus a flat rate model. Cable operators are taking advantage of consumer behavior that shows the faster premium tier service gives them the ability to trade-up existing customers from lower tiers. The availability of a "wideband" tier will amplify this trend, moving the service mix towards a higher average revenue per user (ARPU).

At the same time, as broadband service matures, the types of services included in the

broadband portfolio will grow more complex and bandwidth-intensive. The wideband protocol removes the bandwidth bottlenecks existing today and provides the opportunity for cable operators to re-think their business models in the context of emerging applications. Whether capturing a portion of the music download business with iTunes-like portals, hosting IP-based video download libraries, or rolling out multi-media rich interactive gaming services, cable operators can dramatically expand their service portfolio. In the process, they can significantly increase the average revenue generated per subscriber.

Optimize CapEx Investment

Fully Exploit Today's HFC Plant Potential

A key objective of cable operators today is to find ways to leverage the powerful data capacity of their HFC networks—estimated to use less than two percent of the total available capacity. As seen in the previous section, there are a number of alternatives available, but unfortunately, these are either cost- or labor-intensive (or both).

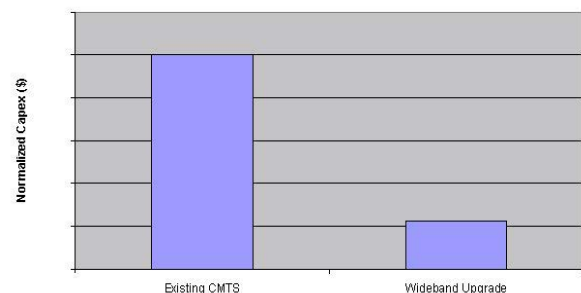
By taking advantage of flexible channel management techniques and channel bonding of up to 24 channels in the first generation of products, the wideband protocol for a DOCSIS network is able to transport packets at Gigabit Ethernet speeds using existing modulation formats—whether they are 64- or 256-QAM. This gives operators the ability to quickly take advantage of the available spectrum on the existing HFC plant, optimizing CapEx investments.

Leverage Existing CMTS Platforms

The flexibility of the wideband protocol for a DOCSIS network does not stop with modulation choices or channel bonding techniques. The protocol can be deployed in parallel with DOCSIS 1.x/2.0 technology, leveraging the investment made in existing CMTSs. Existing edge-QAM modulators can be leveraged as well and physically connected to a wideband module in an existing CMTS. Downstream traffic is now supported from either existing CMTS line cards (DOCSIS 1.x/2.0) or edge-QAM modulators (wideband), and upstream traffic for both variants is supported from DOCSIS 1.x/2.0 line card ports.

The figure below compares the normalized costs for a DOCSIS 1.x/2.0 CMTS with the incremental costs required to upgrade the CMTS to support the wideband protocol. The chart illustrates that with first generation wideband technology, you can more than double the downstream throughput, at less than one fourth the cost! This is the first step towards fulfilling the wideband objective of delivering ten times the throughput at one tenth the cost.

Incremental CapEx



Reduce Downstream Port Cost

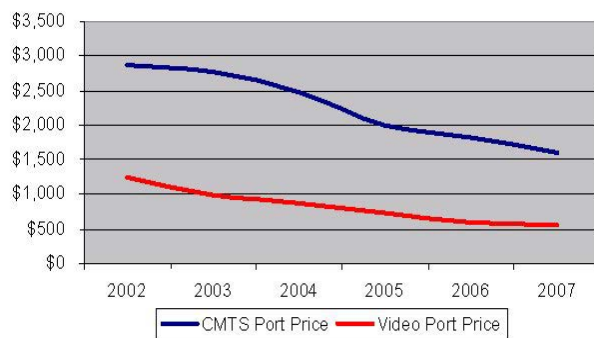
Current CMTS designs are optimized for synchronous traffic for multiple IP-based services. As a result, CMTS line cards are more complex than single-purpose edge-

QAM devices used for uni-directional, asynchronous traffic. This complexity bears with it incremental infrastructure costs.

The proposed wideband technology separates the channel bonding functionality (performed by a wideband module on an existing CMTS) and the Physical Layer adaptation to the RF plant. The latter can be performed by an external edge-QAM device, offering a significant opportunity to leverage declining costs of QAM device designs, and provide a graceful evolution towards a Modular CMTS architecture. In advance of the availability of standards-based products, operators can plan an evolutionary option that saves infrastructure costs.

The following chart depicts the relative costs per port for CMTS and edge-QAM devices. The chart shows that edge QAM devices used in video applications hold a significant price advantage over equivalent CMTS port costs. A CMTS supporting the wideband protocol leverages this advantage by allowing the use of existing QAM devices, thus, lowering the cost of implementation.

CMTS and Video Port Comparison



Source: Infonetics, In-Stat/MDR, Cisco Analysis

Eliminate Port Under-Utilization (Stranded Ports)

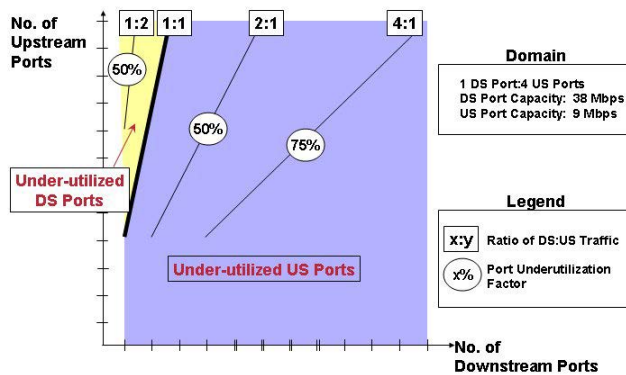
CMTS port assignments were designed around a traffic formula based on studies of anticipated preliminary broadband services and subscriber usage when DOCSIS was

originally formulated. This resulted in the simultaneous support of traffic in upstream and downstream directions and line cards equipped with a fixed ratio of upstream and downstream ports. Typical calculations held to a calculated ratio of one downstream per four or six upstream ports.

The ratio, understandably, was very conservative in the assumptions it made. In addition, it was difficult to predict or anticipate the development of new IP applications that dramatically challenge these assumptions and prevent operators from optimizing port utilization. The net result of this is that some number of ports (either upstream or downstream) will be under-utilized.

The following figure shows how line cards, deployed with a 1:4 port domain are under-utilized for various traffic scenarios. Since the capacity of a downstream channel is roughly four times the capacity of an upstream channel in this example (256 QAM in the downstream, 16 QAM/3.2 MHz in the upstream), the ports in the one downstream – four upstream domain are 100% utilized if the downstream traffic required on the network is equal to the upstream traffic (1:1 traffic ratio line). In the cases where twice as much traffic is required for the downstream than for the upstream (2:1 traffic ratio line), the upstream ports will be 50% underutilized. Finally, if the traffic in the downstream is four times the upstream traffic (4:1 traffic ratio line), the 1:4 port ratio will translate into a 75% upstream port under-utilization.

CMTS Port Utilization



The only way to avoid the under-utilization of upstream or downstream ports is to be able to assign upstream and downstream ports independently of one another. CableLabs' M-CMTS initiative aims at providing this benefit. The wideband technology achieves the same goal with today's CMTS and edge-QAM technology, yielding better utilization of infrastructure and lower cost per subscriber.

Improve Network Efficiency

As a shared access medium, cable networks hold an economic advantage over other access technologies due to the ability to oversubscribe the access plant. Not all users are going to be actively transmitting data at the same time. The determination of an appropriate oversubscription rate is a complex effort—and dependent upon a number of factors such as customer usage patterns, the size and frequency of the traffic bursts, the nature of the content carried, the priority levels for the various types of traffic, and the ratio of peak customer traffic to offered capacity. Cable operators are constantly trying to balance subscriber satisfaction metrics with the need to maintain oversubscription as high as possible in order to have a large number of subscribers sharing the same port, thus reducing the cost per sub.

The wideband protocol offers a simple, yet powerful tool: a “larger pipe”. Complex mathematical models have demonstrated that in the case of bursty traffic, such as the traffic generated by Internet users, a higher capacity transmission medium will show better statistical multiplexing characteristics since the probability of transmission collisions are reduced. Wideband offers such statistical multiplexing gain by logically bonding multiple QAM channels and increasing the total channel capacity. In other terms, this means that logically bonding multiple QAM channels will increase the number of subscribers per QAM (or per port).

Looking ahead, both higher data rate tiers, as well as emerging applications such as IP-based video will change the traffic characteristic. In this scenario, the number of subscribers that will be able to share a single channel will be significantly reduced, making the current technology less and less profitable. The wideband protocol offers the ability to maintain the same level of oversubscription rates, despite increasing customer demands for higher throughput.

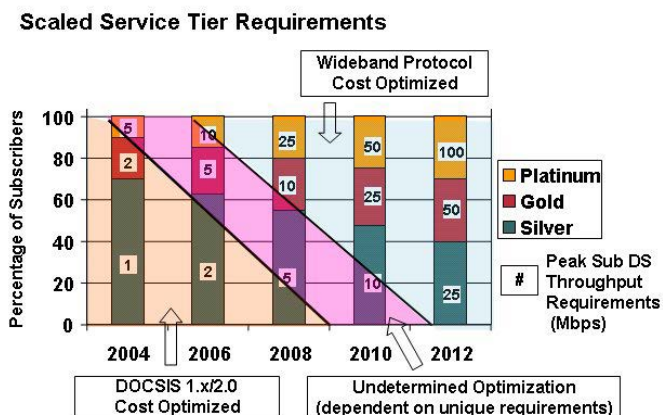
Scale to Future Requirements

Wideband technology offers deployment flexibility. Because a wideband logical channel can be dynamically defined to be a bundle of two, four or any number of QAM channels up to 24, as defined in my initial proposal for DOCSIS 3.0, cable operators will be able to choose, over time, the required channel capacity according to their business needs.

On the modem side, cable modem tuner technology is available today to allow multiple QAM channels to be received simultaneously, enabling wideband service. Such technology represents an important evolution because it allows multiple channels to be demodulated by a single, digital

multichannel receiver chip, instead of multiple discrete traditional receivers, adding complexity and cost to the modem. This means that first generation wideband modems will have built-in capability to receive up to 16 channels, enabling the operator to scale to future bandwidth and throughput requirements, without changing the wideband installed base.

The figure below illustrates projected high-speed data services that cable operators may see in the coming years. In each case, the service offering is split into three or more tiers. A graphic (the diagonal stripe) is superimposed on the chart, indicating the long-term costs of using wideband protocol versus DOCSIS 1.x/2.0 technology. The precise placement of this diagonal stripe is determined by exact costs for the CMTS and cable modem expenses, as well as oversubscription rates.



Regardless, the conclusion is that for many years to come, the most cost-effective network is one that simultaneously delivers services to both sets of customers.

Minimize Operational Cost

Backwards-Compatible with DOCSIS 1.x/2.0

Cable operators must consider how to continue to innovate and introduce new services to the market, without abandoning

their current customer base or requiring costly equipment upgrades. The wideband protocol for a DOCSIS network addresses this by leaving unchanged the DOCSIS 1.x/2.0 network and the subscribers it supports. Yet it provides incremental opportunities to begin servicing customers that demand higher data rates.

Looking ahead, cable operators will increasingly support service tiers and packages that offer varying throughput speeds and quality assurances. This will inevitably fragment customer requirements; many will focus on lower-speed, value-priced services, while a smaller percentage will demand higher-speed services. The challenge to cable operators will be to cost-effectively serve all customers.

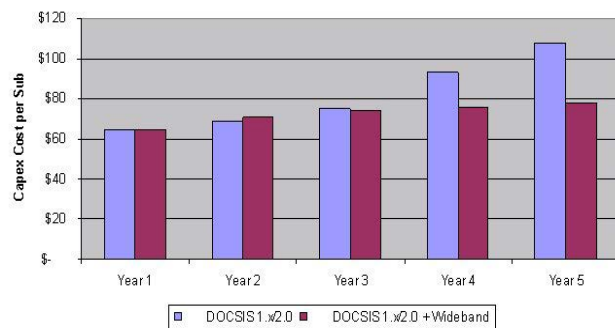
The wideband protocol is ideally suited for this situation. Because it builds upon existing CMTSs and the DOCSIS 1.x/2.0 protocol, it does not affect the customers that are currently served on the plant today. New customers can be served simultaneously with DOCSIS 1.x/2.0 customers from the same CMTS. Additionally, existing provisioning systems and operational processes can be leveraged. This offers an ideal environment for cable operators to employ, given the flexibility that is offered.

WIDEBAND BUSINESS CASE SUMMARY

The previous sections have detailed a multitude of factors that each contribute to the business case for wideband technology. This section illustrates the combined effect of each of these factors in providing a compelling business case for deploying the wideband protocol. The charts that follow summarize the financial results for two deployment scenarios: one using DOCSIS 1.x/2.0 technology, and one using both DOCSIS 1.x/2.0 and wideband. The model is based upon a cable footprint of 1 million HHP, with

existing high-speed data take rates of 25% in year 0 (to reflect the embedded base of equipment and customers) increasing up to 40% in year 5. The chart illustrates how the wideband technology significantly lowers CapEx per subscriber.

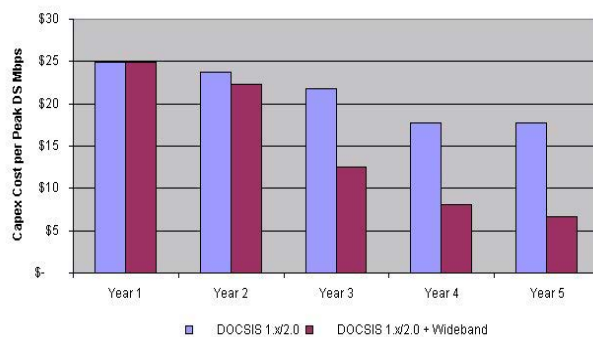
CapEx Cost Comparison – Same Service Mix



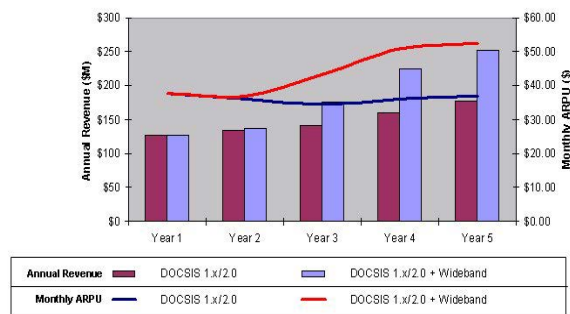
But this direct comparison tells only half of the story. A primary reason for using wideband technology is to offer even higher service throughputs, delivering more value, and capturing more revenue. The remaining charts summarize the financial benefits when the wideband protocol is deployed, enabling cable operators to deliver services as high as 100 Mbps data throughput.

The CapEx cost per Mb will decrease significantly as a result of using the wideband protocol. The improved revenue and earning results speak volumes to the benefits of being able to offer higher-speed services.

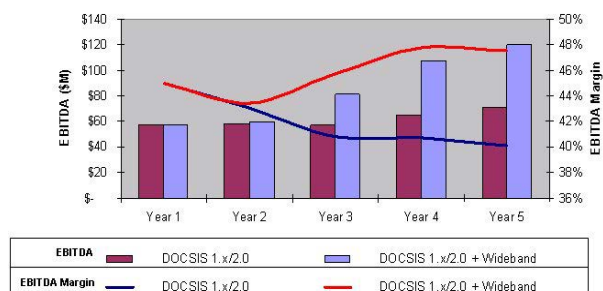
CapEx Cost Comparison – Different Service Mix



Revenue Comparison – Different Service Mix



Earnings Comparison (EBITDA) – Different Service Mix



The following conclusions can be made from this business analysis:

- Wideband can dramatically lower network CapEx as subscriber throughput rates increase
- CapEx incremental costs required to support the wideband protocol are modest and significantly drive down per Mb cost
- Wideband offers the ability to retain and continually upsell existing customers to higher service tiers
- Wideband offers the ability to deliver higher-value and revenue
- Operating margins can be improved by offering higher priced service tiers

This analysis captures only the cost of delivering data capacity. Incremental services delivered over this transport offer further upside, particularly in the case of higher data rate services

FULFILLING CABLE'S VISION

Cable operators were the first, and arguably the most credible service providers

to articulate the vision of an intelligent, flexible, secure, scalable network suitable for supporting multiple services simultaneously. The use of IP, and the framework established by DOCSIS have been key factors leading to the articulation and fulfillment of that vision. But the vision can only be fulfilled if the full capacity of the network is leveraged and made available to the end user.

The wideband protocol for a DOCSIS network, with its ability to unleash the full potential of the HFC network, is the key to making this happen. No longer will cable operators be blocked by artificial restrictions imposed by historical assumptions and RF channel management techniques. For as much of the spectrum an operator allocates towards IP-based services, a customer can theoretically access. But while the wideband protocol for a DOCSIS network can provide access to this inherent competitive advantage, the real challenge will come in how operators translate this potential into a true strategic business initiative.

Cable won the initial battle of broadband and the triple play. Will the industry be able to win the next round? Wideband and subsequent strategies lay the foundation for this and provide the catalyst for change for years to come. The outlook for the cable industry's success very much depends upon how effectively cable operators leverage this advantage. But there is an additional benefit that is as important if not more so. While the wideband protocol for a DOCSIS network makes more efficient use of the bandwidth inherent in the HFC network, it also gives operators the ability to unleash the full power of IP, as well as take a major step towards achieving true network and service convergence. As the subject of follow-on papers, this will be the remaining piece to our challenge in achieving a successful business plan that embraces a true wideband perspective.

REFERENCES

- Chapman, J., "The Wideband Protocol for a DOCSIS Network," 2005 SCTE Conference on Emerging Technologies
- Eng, J., "Next Gen Full-Service, All-Digital HFC Network Beyond DOCSIS 2.0", 2004 SCTE Conference on Emerging Technologies

ABOUT THE AUTHOR

John T. Chapman is currently a Distinguished Engineer and the Chief Architect for the Cable Business Unit at Cisco Systems in San Jose, California. As a founding member of the Cisco Cable BU, John has made significant contributions to Cisco and the cable industry through his pioneering work in DOCSIS and development of key technologies and concepts critical to the deployment of IP services over HFC plants.

Included in these achievements are being the primary author of significant portions of the DOCSIS and PacketCable specifications as well as the originator of DOCSIS Set-top Gateway (DSG) and evolving specifications for DOCSIS Wideband and Modular CMTS architectures for the industry's Next Generation Network Architecture (NGNA) initiative. John has also published a number of ground breaking whitepapers on Multimedia Traffic Engineering (MMTE), DSG, QoS, and high availability and is a respected and frequently requested speaker at industry events.

John has 18 patents issued and 27 patents pending in a variety of technologies including telephony, VoIP, wide area networking, and broadband access for HFC cable networks. In his spare time, John enjoys spending time with his wife and two daughters. John is a 6th Degree Black Belt Master in Tae Kwon Do and enjoys white water canoeing and skiing.

Previous papers by John may be found at <http://www.johntchapman.com>

MAXIMIZING BANDWIDTH UTILIZATION VIA ADVANCED SPECTRUM MANAGEMENT

Jack Moran

Distinguished Member of the Technical Staff
Motorola Connected Home Solutions

Abstract

This paper discusses the impairments that must be addressed to maximize bandwidth utilization over DOCSIS infrastructure. It presents an effective approach for overcoming these impairments and delivering increased throughput for DOCSIS 1.X and 2.0 cable modems by taking advantage of DOCSIS 2.0 extensions. It also shows how extensions to DOCSIS can allow operators to implement advanced spectrum management to cost-effectively buildout standards-based infrastructure today while retaining the flexibility to support emerging or evolving architectures in the future.

INTRODUCTION

The twin demands of more cost-effectively utilizing existing buildouts while delivering increased capacity to the customer require innovative means of maximizing the utilization of existing bandwidth. Clever new approaches are required to optimize the use of existing cable modems while achieving increased throughput for DOCSIS 2.0 modems.

Advanced spectrum management can be deployed along with DOCSIS standards to allow operators to improve upstream throughput by dynamically evaluating the RF characteristics of available upstream spectrum and then selecting the spectrum frequency, modulation mode (16 QAM, 32 QAM, 64 QAM, 128 TCM, 256 QAM) and channel width (1.6 MHz through 6.4 MHz) to optimize throughput.

Extensions beyond DOCSIS 2.0 can allow operators to monitor RF performance in a non-intrusive manner to truly understand the impact of impairments.

Once these impairments are carefully measured, operators will then possess the information they need to operate at the highest throughput possible at all times. This is enabled by the ability to implement sophisticated ingress noise cancellation and impulse noise immunity techniques. With the majority of cable modems installed today supporting only DOCSIS 1.0, there are also advanced spectrum management extensions available to allow operators to nearly double the throughput of their installed basis of DOCSIS 1.0 modems by enabling them to run 16 QAM modulation virtually anywhere that QPSK is currently running.

MEASURING SPECTRUM

The first step toward managing spectrum is to measure and understand the impairments on the DOCSIS infrastructure.

Without a clear understanding of the impacts of impairments, it is difficult to optimize throughput and meet or exceed DOCSIS performance goals.

What is truly required from a DOCSIS Cable Modem Termination System (CMTS) is the ability to intelligently assess the impairments on a given return path in real time. Unfortunately the time required to

make an accurate assessment of impairments using traditional measurement tools far exceeds the time budget allotted for measurements.

Simply stated, the more time that is required to perform an accurate measurement, the more the throughput of the data traffic is impacted.

While Fast Fourier Transform (FFT) measurements are a useful analysis tool for ingress and impulse noises, this technique has little ability to accurately measure micro-reflections and group delay. Instead, coherent measurements are required to measure the effects of linear distortions such as micro-reflections, amplitude, and group delay distortion.

When discussing the system non-linearity class of impairments, only a coherent measurement can be used. And, the larger the QAM constellation, the greater susceptibility to system non-linearity. This has been demonstrated by comparing the difference between 16-QAM and 64-QAM given the same non-linear circuit and the same RMS power for both constellations.

The entire point of advanced spectrum management is to assess the unused return path bandwidth for all impairments. This takes time an active data channel simply cannot afford to spend.

ADVANCED SPECTRUM MANAGEMENT DEFINED

Advanced spectrum management allows operators to identify impairments and make the necessary adjustments to improve performance. An effective method to perform advanced spectrum management is to utilize a dedicated receiver on the CMTS to monitor performance on any one of the upstream paths without impacting performance.

Advanced spectrum management is an extension to DOCSIS that essentially relies on a spare receiver to perform time-consuming measurements in the background.

Operators can therefore gain access to all of the return nodes connected to one of the receiver ports and perform tests on any available modem on any one of the receiver port's supported nodes. Advanced spectrum management is a necessity for any cable operator to be able to efficiently support VoIP or other demanding real-time services.

Given the unknown limitations that exist in the return path, the DOCSIS CMTS must assist the cable operator in determining what any given return path is capable of yielding. Advanced spectrum management must be completely transparent, with absolutely no impact on voice, data, or video throughput. It must also be able to discern between a linear and non-linear distortion to be truly effective.

Advanced spectrum management is critical to the ability to optimize the billable capacity of the DOCSIS network. Operators simply cannot fix impairments that they cannot identify and measure. Hard data is necessary so that the effects of multiple impairments can be discerned and successfully addressed.

The only alternatives to advanced spectrum management are to ignore (or guess at) the impacts of impairments, or to deploy expensive, dedicated testing gear to measure the various impairments that are present on any given DOCSIS network. Operators have relied on advanced vector signal analyzers and next generation CATV Analyzers that support spectrum and DEMOD measurements for several years to characterize the return path characteristics.

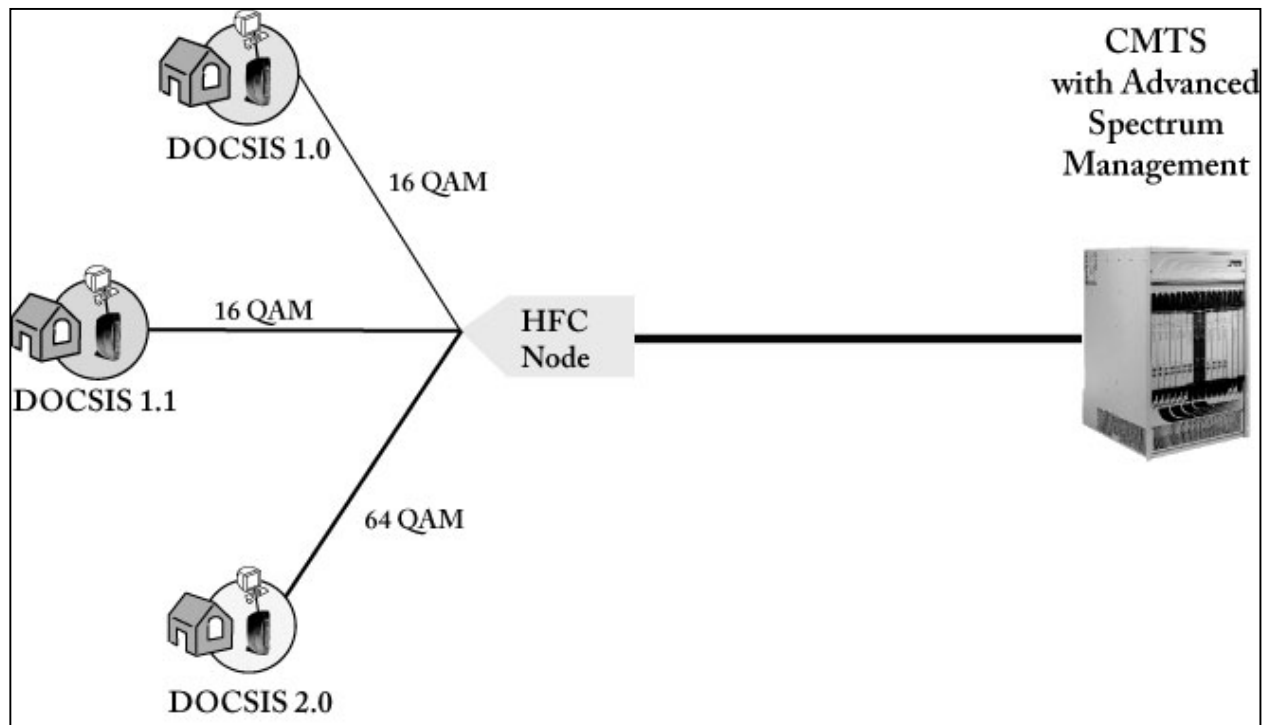
But this approach is expensive, difficult to deploy and extremely time-consuming to measure. It also impacts the performance of the DOCSIS network, requiring operators to take a segment of the network out of service as the reference sources tend to be constant carrier signals that disrupt service in some instances.

IMPLEMENTING ADVANCED SPECTRUM MANAGEMENT

By implementing advanced spectrum management techniques in conjunction with the DOCSIS specifications, operators can ensure bandwidth-efficient co-existence between DOCSIS 1.0, 1.1, and 2.0 cable modems. They can deliver higher upstream bandwidth that can even exceed DOCSIS specifications while ensuring a successful transition to DOCSIS 2.0. Operators can

The receiver technology provides post-equalization support that can double the throughput for existing customers because virtually all DOCSIS 1.0 cable modems able to run in QPSK would be able to operate in 16 QAM mode utilizing a post equalization technique. When there is a significant number of DOCSIS 2.0 cable modems installed, the cable operator can begin the ATDMA Logical Channel Operation in which the Symbol Rate remains the same (2560 ksym/s) but the DOCSIS 2.0 cable modems can begin to transmit in a pure ATDMA mode of operation, i.e. with extended Forward Error Correction, byte interleaving (if necessary), and higher constellation rates such as 32 QAM or even up to 256 QAM.

The financial benefits of this migration approach are compelling. Operators can



transition to 2.0 without providing a performance burden to legacy subscribers because the CMTS can operate in DOCSIS 1.X mode.

accelerate revenue from 2.0 services, and they can implement gradual migration at the pace that makes the most economic sense for them. They can continue to support legacy

modems while introducing new services to these subscribers, and they can concurrently support DOCSIS 1.X and 2.0 operation across the same infrastructure.

Cable operators can double the upstream bandwidth for a large population of modems, thus creating increased billable bandwidth without further network buildout. They can create upstream bandwidth that supports higher-speed services and enable new broadband services that command premium pricing. This approach is not only the most practical migration path; it is also the one with the lowest risk to the cable operator.

To take advantage of these business opportunities, operators first need to address the transient impairment issues that today constrain upstream bandwidth in the return path.

MEASURING IMPAIRMENTS

Operators need to be able to improve the Signal-to-Noise Ratio (SNR) so they can more efficiently manage spectrum and improve noise cancellation simultaneously over diverse populations of DOCSIS 1.0, 1.1, and 2.0 modems. Advanced spectrum management is important so that operators can understand the various impairments present in their infrastructure.

This is particularly critical because real-world environments face the following three major classes of impairment, which are present at some level the majority of the time:

- Linear Impairments
- Non-linear Impairments
- Transient Impairments

LINEAR IMPAIRMENTS

Micro-reflections are the most common distortions that exist in every plant, but they occur differently in each plant. They are caused by impedance mismatches, and the most significant micro-reflections tend to occur at the lower tap values of the coax plant.

The lower the tap value, the poorer the isolation between the other ports and the cable modem signal. Micro-reflections are frequency dependent, so that not all channel bandwidth is affected equally. A fundamental problem in isolating a micro-reflection lies in the fact the impedance mismatch resulting in a rather large micro-reflection tends to be dominated by a poor termination in a tap adjacent to the problem cable modem and not to the tap that the cable modem is connected to. This phenomenon is directly related to the poor isolation characteristics of the low tap values.

Amplitude tilt or slope distortion is also present in every cable plant, and can be caused by coaxial cable loss or more likely by the use of diplex filters. Every plant also faces group delay—or “phase” distortion—which becomes a bigger problem as filtering is introduced. The major source of group delay distortion is diplex filters. The more amplifiers in cascade, the more dramatic the impact of both amplitude and group delay distortion on a DOCSIS transmission.

Advanced spectrum management allows operators to address these types of linear impairments by:

- Migrating away from the problem frequencies
- Reducing the symbol rate, usually by half
- Equalizing the distortion

NON-LINEAR IMPAIRMENTS

There are no specifications that exist in DOCSIS 2.0 that address system non-linearity. DOCSIS 2.0 Technology—particularly any modulation type greater than 16-QAM—is inherently weaker against non-linearity than is DOCSIS 1.X technology.

This is due to the higher crest factor of DOCSIS 2.0 that ranges from a minimum of 3 dB higher to a maximum of 7 dB higher peak power. There are many variables that contribute to the final crest factor, such as the:

- Theoretical crest factor of the modulation constellation
- Number of diplex filters in cascade
- Carrier frequency

By definition non-linearity is a signal-dependent distortion, which simply means that the effects of a non-linearity can only be observed in the presence of a signal.

Non-linear impairments include Common Path Distortion (CPD) and not-so-common path distortion, often referred to as return laser non-linearity.

CPD is well understood by the CATV industry in general. It is the phenomena of a coaxial connector becoming or temporarily acting as a diode. It is easily observed by seeing analog video carriers spacing 6 MHz apart throughout the return path. While CPD is easily detectable, the return laser being either clipped or just becoming marginally non-linear can only be witnessed today by advanced spectrum management on a dedicated receiver on a CMTS card or by deploying vector signal analyzer test equipment on the network.

With advanced spectrum management, one can easily observe that the effects of any non-linearity is that the outer constellation points are impacted far greater than the inner constellation points.

TRANSIENT IMPAIRMENTS

Transient impairments include ingress noise and impulse noise. Narrowband AM modulation carriers such as shortwave radio signals can suddenly appear anywhere in the return path spectrum. Ingress noise refers to any interference that is coupled into the return path plant via an external source.

The predominant coupling mechanism for ingress noise is a poorly shielded drop coaxial cable that is acting more like an antenna than a drop cable. The overwhelming majority of ingress noise is narrowband AM modulated carriers whose bandwidth is usually less than 20 kHz and ingress noise in general seldom has a bandwidth over 200 kHz. Return path characterizations conducted by Motorola over time have found that this ingress interference ranges from around -25 dBc (25 dB below the DOCSIS signal power) to +15 dBc (15 dB above the DOCSIS signal power).

Occasionally, ingress noise is recorded that is as high as +25 dBc. It is therefore important to ensure that the CMTS receiver front end can take at least +31 dBc or 6 dB more than the +25 dBc ingress incidents. Frequency avoidance has been the sole technique to deal with this parameter until the development of advanced ingress noise cancellation techniques.

Impulse noise is also a reality in virtually all return path cable plants. It is made up of short bursts of high-level noise such as that resulting from the coupling of transients into

a channel. Wideband noise events occur typically in bands wider than 6 MHz, and there are multiple sources of this type of impairment.

However, the duration of this noise usually lasts in the 1-100 microsecond range. There is another class of impulse noise that is powerline related, and when this type of transient event occurs, the duration is in the 1-10 millisecond range.

The one saving grace of either type of impulse noise is that neither has any significant energy beyond 15 MHz. Evidence of this fact is that virtually all DOCSIS 1.X systems operate error free even in 16 QAM modulation with a moderate amount of Forward Error Correction enabled when operating over 20 MHz.

UNDERSTANDING PERFORMANCE

Advanced spectrum management allows tremendous flexibility for understanding the complex interaction of impairments on the DOCSIS network. No single impairment can be clearly singled out for testing because most—if not all—impairments are present at some level the majority of the time. It therefore becomes a matter of assessing the magnitude of each impairment's impact on the DOCSIS service.

For example, conducting an impulse noise performance test without also measuring ingress noise performance is not particularly insightful. Another fundamental challenge is that measurement time directly impacts throughput, because in the typical scenario you cannot send data while you are taking measurements.

Operators trying to improve performance are hard-pressed to impose increased demands on the infrastructure by performing

continuous testing that negatively impacts the bandwidth being tested. Hence the paradox; many CMTS platforms are only performing the nearly transparent FFT measurement because operators cannot afford to impose the overhead of measurement.

Many operators therefore cannot afford the time for a coherent measurement approach—even though it is universally agreed that coherent measurement is the only accurate assessment of any impairment's impact on a DOCSIS service.

Advanced spectrum management implemented on a CMTS with spare receiver ports on interface cards can monitor performance on any one of the upstream ports without impacting performance. It can non-obtrusively gain access to all of the return nodes connected to one of the receiver ports and perform tests on any available modem on any one of the receiver port's supported nodes. The spare receiver is effectively connected in parallel with a selected receiver port so the operator can measure traffic and performance in real-time on any given live receiver port.

It should have access to all of the mapping information as well as a full list of cable modems available to whichever receiver port is currently being evaluated. Therefore, while the receiver port being monitored is performing its function at full capacity, the spare receiver has the luxury of time to perform detailed, lengthy, and coherent SNR measurements. It can also perform a host of other measurements by simply borrowing an idle cable modem for a rich set of return path calculations. The borrowed cable modem is automatically released for service if demands are placed upon it and another idle modem is selected to ensure there is no intrusion on customer service.

Advanced spectrum management also requires support for adaptive noise cancellation at the receiver to measure the diverse types of noise, process this information and take action to cancel it out in real time. For example, an ingress canceller on the CMTS can track and cancel rapidly changing severe CPDs. The net effect is that the operator is able to maintain a high-order QAM modulated digital carrier.

If the noise cannot be cancelled out—such as a very large ingress noise or interferer—the CMTS can avoid the noise by changing the modulation mode or moving frequencies. Operators can therefore continuously improve performance, proactively recognize and resolve potential bottlenecks, and create more billable bandwidth.

Continuous monitoring and adaptation allows cable operators to aggressively implement advanced noise cancellation in environments where the types and degrees of noise change frequently. These impairments have historically been the limiting factors in achieving QAM modulation higher than 4 QAM (QPSK). The combination of post equalization and superior ingress noise cancellation capabilities results in a DOCSIS 1.X system today where 16 QAM, error-free operation is achievable virtually anywhere in the return path.

SAMPLE MEASUREMENTS

The following is just a small sampling of the types of measurements operators can implement using advanced spectrum management:

- SNR per cable modem with post equalization enabled.
- SNR per cable modem with post equalization disabled, which reveals whether there is any significant micro reflection on a per-modem basis.

- SNR per cable modem with ingress noise canceller enabled.
- SNR per cable modem with ingress noise canceller disabled, which reveals whether there is any significant ingress noise present.
- RX level per cable modem, which when coupled with the cable modem transmit level from the SNMP MIB will allow operators to calculate network loss per cable modem.

The reality is that there is a virtually unlimited number of tests operators can perform using advanced spectrum management without impacting the performance of active cable modems.

CONCLUSION

It is very difficult to improve something that you cannot measure. Operators need hard data on network performance, but cannot afford the cost and complexity of deploying test equipment throughout the network.

Merely guessing at the impact of impairments is a wasteful and frustrating exercise. The ability to accurately measure and monitor impairments is essential to optimizing the productive use of bandwidth and ensuring the successful delivery of real-time video and voice services.

By implementing advanced spectrum management, operators can non-obtrusively monitor the impact of impairments on an ongoing basis and take corrective actions when necessary to improve performance of voice, data, and video services.

ABOUT THE AUTHOR

Jack Moran is a Distinguished Member of the Technical Staff for the IP Solutions

Group for Motorola Connected Home Solutions. He is the holder of 11 U.S. patents in data communications, with many more pending.

Moran is responsible for DOCSIS Physical Layer performance over an HFC RF network system. For over four years, he has been modeling the return path for DOCSIS 1.0, 1.1, and now 2.0 performance capabilities. This modeling effort has

included many live plant characterizations in an effort to simulate on a repeatable basis the type of real-world impairments that DOCSIS systems must overcome.

He is also a member of the DOCSIS 2.0 Technical Team PHY Layer Issue as well as a member of the former IEEE 802.14 Cable Modem Group. Moran can be reached at jack.moran@motorola.com.

MAXIMUM CAPACITY: THE ROLE OF INTELLIGENT EDGE DEVICES IN CABLE NETWORK CONVERGENCE

Michael Adams
Terayon Communication Systems

Abstract

The ultimate goal of many cable systems operators is the migration to an IP end-to-end delivery model. However, there are many constraints at the HFC level such as existing analog channel assignments, must-carry regulation, legacy set-top box investment, and interface standardization.

Nevertheless, the internal structure of cable systems is evolving rapidly towards a converged IP backbone that carries all content distribution and signaling.

Therefore what is needed is an intelligent edge device that can perform format conversion, protocol translation, and content localization so that the IP backbone can be implemented today, with the ultimate goal of extending IP carriage all the way to the home.

INTRODUCTION

Today cable operators are typically managing four, discrete network protocols over their HFC plant:

- Analog broadcast video
- Digital broadcast video
- On-demand video
- High-speed data and VoIP

The HFC bandwidth allocation for each is statically defined as 6 MHz channels. It is difficult to re-distribute network capacity between services because each has a separate

management system, which doesn't 'know' about the other services.

Most cable operators also manage two or more transport networks (regional, and headend-to-hub interconnections). The history of technical development is often the driver for the right choice at the appropriate time, and four different technologies are in common use today:

- AM Supertrunking
- IP over ATM over SONET
- IP (Packet) over SONET
- IP over Gigabit Ethernet

The complexity of protocol conversion and bandwidth adaptation at the edge is significant, but it is now possible to build cost-effective, intelligent edge devices that allow a single transport network and considerable simplification of the HFC network. These developments also support efficient bandwidth management and a graceful migration to a single, converged IP network topology, which will enable rapid deployment of new services without having to invent new operation management systems.

This paper will describe this evolution in detail, and show how flexibility and cost-effective deployment can be achieved at each step of the process. The steps described will cover:

1. Compression of all services to a common coding and transport standard.

2. Transmission of all services over a single, IP-based transport network.
3. Localization of video services to the zone or set-top level using digital program insertion (DPI). Switched Broadcast services may leverage the same infrastructure that is used for localization.
4. Adaptation of services to maintain backward compatibility with legacy devices in customers' homes. (This is especially important for regulatory reasons and to preserve the multi-billion dollar investment the industry has made in MPEG-2 set-tops.)

HEADEND AGGREGATION AND GROOMING

Cable networks have developed over the years by adding new services in an incremental fashion – each new service bringing with it new headend equipment and leveraging existing network transport. As a

result of this the cable operator is faced with the challenge of operating multiple different systems, each with its own operational quirks.

For example, when digital services were added, the most cost effective deployment was to modulate and combine the analog and digital channels at the headend and to use existing AM supertrunking to deliver that combined signal to the distribution hubs. Figure 1 illustrates the basic transformations that are done at the headend to allow video feeds to be selectively groomed by the operator into an optimal channel lineup.

As shown in the diagram, all functions that transform the content are done at the master headend. From the master headend, the network uses Amplitude Modulation (AM) Supertrunking to distribute the RF signals to the hubs and ultimately to the viewer's home.

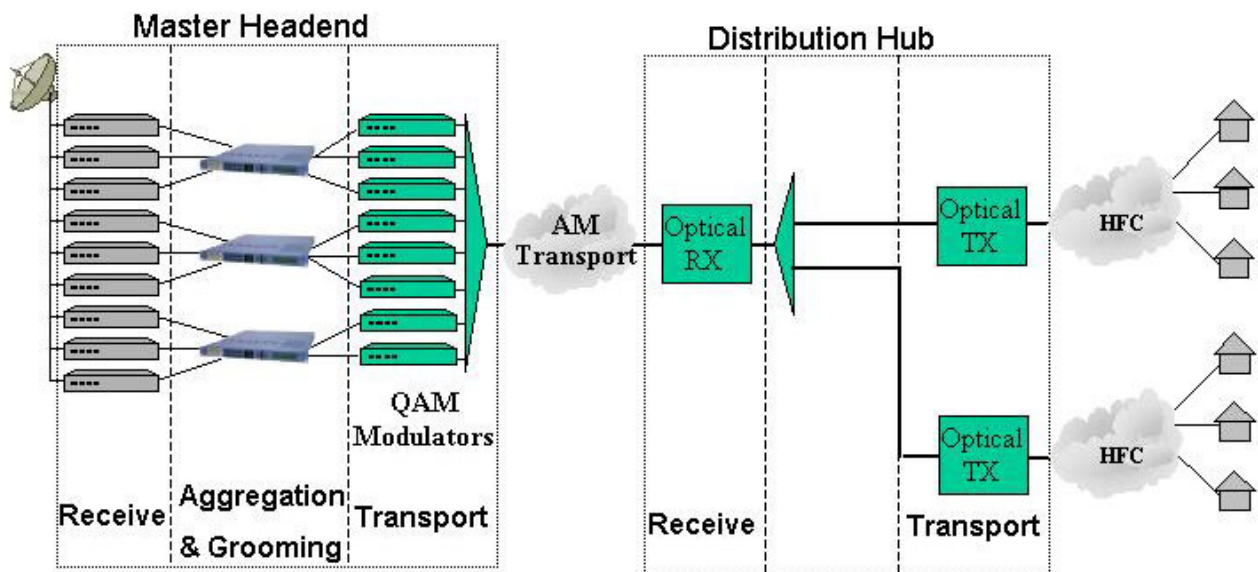


Figure 1: Headend Aggregation and Grooming

In Figure 1, all digital feeds from satellite are fed into a series of groomers, which perform a number of related functions:

1. The feeds are aggregated together, allowing two satellite feeds to be typically combined into a single QAM channel over the cable system.

2. The feeds are groomed to remove any programs that are not required for carriage over the cable system
3. The feeds are re-mapped so that any Packet Identifier conflicts are removed and program map and program association tables built for the consolidated transport stream output.
4. The programs are statistically multiplexed so that the instantaneous bit rate does not exceed the output channel limitation (38.8 Mbps for a 256-QAM channel).

Although this approach has been extremely successful in allowing cable operators to design their own efficient digital channel line-ups, there are some significant disadvantages:

1. The digital line-up is fixed for all parts of the cable system – this may be a problem if different parts of the system are not upgraded to the same capacity.
2. As regional clusters are interconnected, it is possible to push much of the headend functionality back to a regional super-headend.
3. The AM supertrunk is still being used to deliver the digital video channels. As cost-effective digital transport technologies are becoming readily available and are being used to Video On Demand (see later). It is becoming practical to unify all transport over a single, IP transport system.

For these reasons, an all-digital distribution model is rapidly gaining acceptance. As we will see, it is also a foundation for Digital Program Insertion, Video-on-Demand, and Digital Simulcast.

Of course, all digital distribution solves a number of problems. There is adequate capacity available in modern 10 Gbps multi-wavelength transport systems to carry all

services over a single backbone. The single biggest issue becomes how to carry feeds over the backbone that are only available in analog format. Until recently it was prohibitively expensive to MPEG encode analog channels, but as encoding costs have come down, and as cable systems have grown larger (so that the cost of encoding an analog channel is spread over a larger subscriber base) this is no longer such an issue.

The other issue, again until recently, was the need to do local commercial insertion into analog channels. This has been solved by the implementation of digital-in-digital insertion systems that operate at the MPEG layer, using a standardized approach called Digital Program Insertion.

DIGITAL PROGRAM INSERTION

The basic principle of Digital Program Insertion (DPI) is that at a given digital cue signal (signaled using the SCTE 035¹ message) an individual output program can be seamlessly spliced from the network feed to a local feed generated by a server. The DPI system has been divided into two main sub-systems; the Ad Server and the Splicer. The two communicate using a set of standard messages according to SCTE 030².

After the local commercial is inserted, the program is spliced back into the network feed. To do this successfully, the splicer has to be aware of the MPEG decode buffer model and the structure of the MPEG encoding syntax. In practice the bit rate of the inserted commercial has to be modified to ensure a smooth transition without frame drop or repeat, therefore rate-shaping technology is incorporated into the DPI splicer.

Early deployments of DPI were implemented at the MPEG-2 physical layer

using Asynchronous Serial Interface (ASI) interconnections. However, DPI is becoming the first step in the migration to an IP network layer. To achieve this the MPEG-2

transport packets are encapsulated into a UDP flow, which can be carried over an IP network.

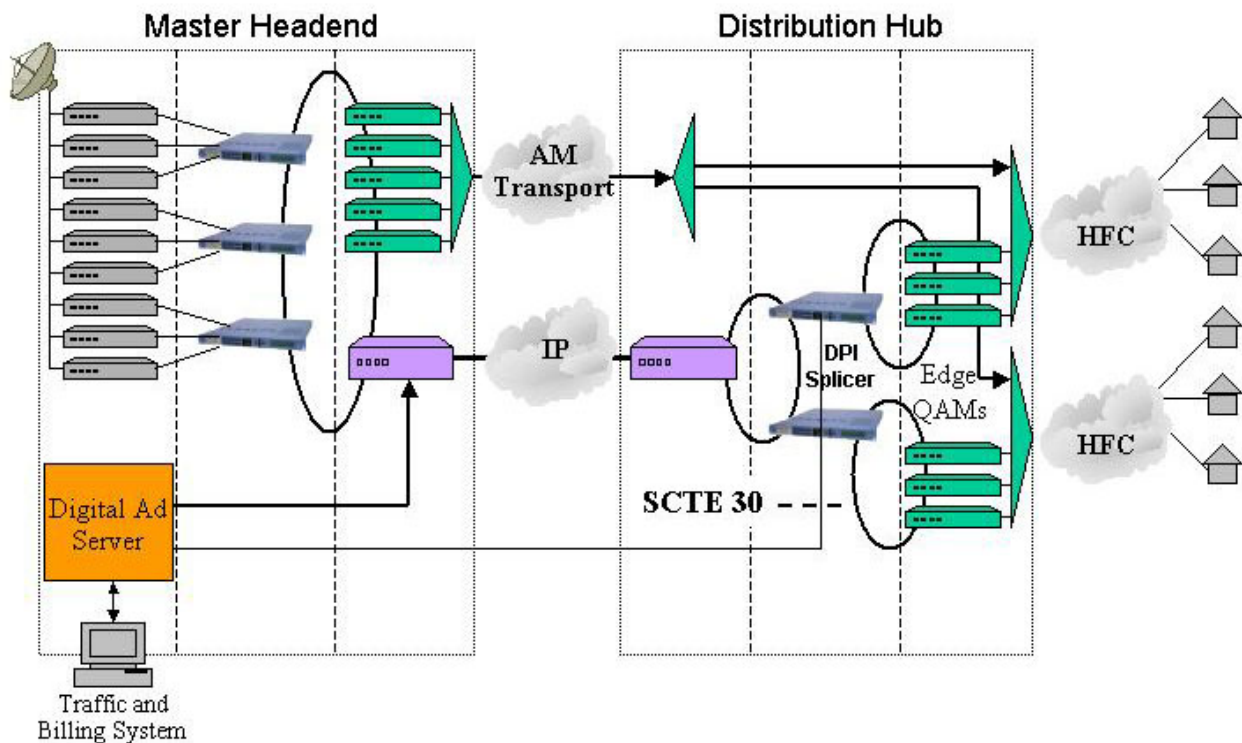


Figure 2: Digital Program Insertion

Figure 2 shows the latest generation of DPI implementation. There are a number of key points:

1. The Digital Ad Server does not have to be co-located with the DPI splicer because they are connected via an IP connection (and not a distance-limited ASI connection).
2. The Digital Ad Server can be centralized at the headend for ease of operations and maintenance.
3. The DPI Splice can be located at the edge of the backbone network at the Distribution Hub. This allows the operator to provide zoned ad insertion to an arbitrarily small serving group area.

DPI is a sophisticated new technology with tremendous flexibility. A good

introduction to the DPI standard is available in SCTE 067³.

VIDEO ON DEMAND

Video-on-Demand is being rolled out aggressively by all major MSOs. The technology has evolved to larger capacity servers with Gigabit Ethernet or 10 Gigabit Ethernet output ports connected over an optical IP transport network to high-density edge-QAM devices.

Figure 3 shows a typical current implementation of VOD. Note that the VOD server is centralized and that the Edge QAMs are distributed to the distribution hub sites. In this diagram a potential future implementation of rate shaping is shown, in which CBR streams are statistically

multiplexed before delivery to the edge-QAMs. In the current CBR implementation, 10 streams encoded at 3.75 Mbps per stream can be delivered over a 256-QAM channel. Depending on content, a significant gain of maybe 2-4 streams per QAM can readily be achieved. In some cases this technology may be used to accommodate peak demand without blocking.

The important thing to note is that Figure 3 is identical architecturally to Figure 2. This

is a nice property as the same equipment can be easily re-configured between broadcast, DPI, and VOD services. In addition, Switched Digital Broadcast can also be accommodated using this same basic architecture, the additional complexity of Switched Digital Broadcast being extensions to the control plane to allow set-tops to request channels (in the same way that VOD allows a set-top to request sessions).

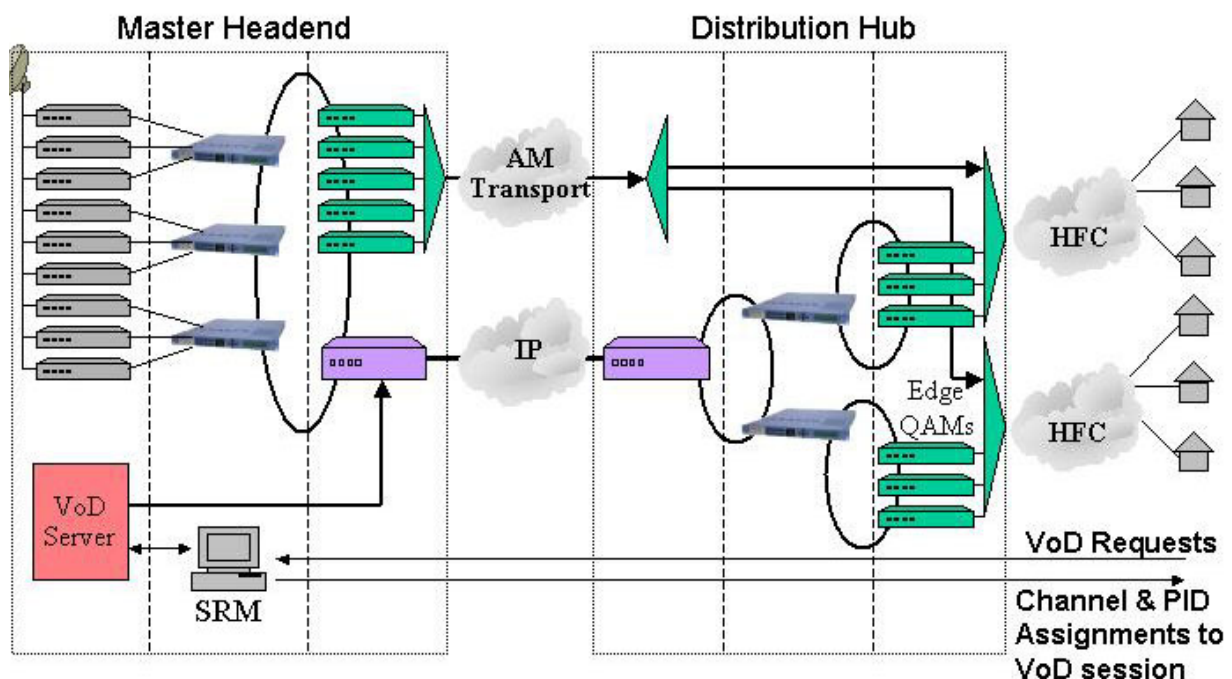


Figure 3: Video-on-Demand

DIGITAL SIMULCAST

Digital Simulcast makes all programming available as part of a digital tier. To do this any analog feeds must be encoded, even off-air signals. Figure 4 shows the processing at the headend to support digital simulcast. Once all of the program feeds are in MPEG-2 format, they can be encapsulated into IP packets for transfer over the converged IP

backbone network, which also supports all other connectivity required in the system for on-demand programming, high-speed data, and voice.

When the signals arrive at the distribution hub, an intelligent edge device terminates them and distributes them to the various legacy channels for distribution over the HFC plant as shown in Figure 4.

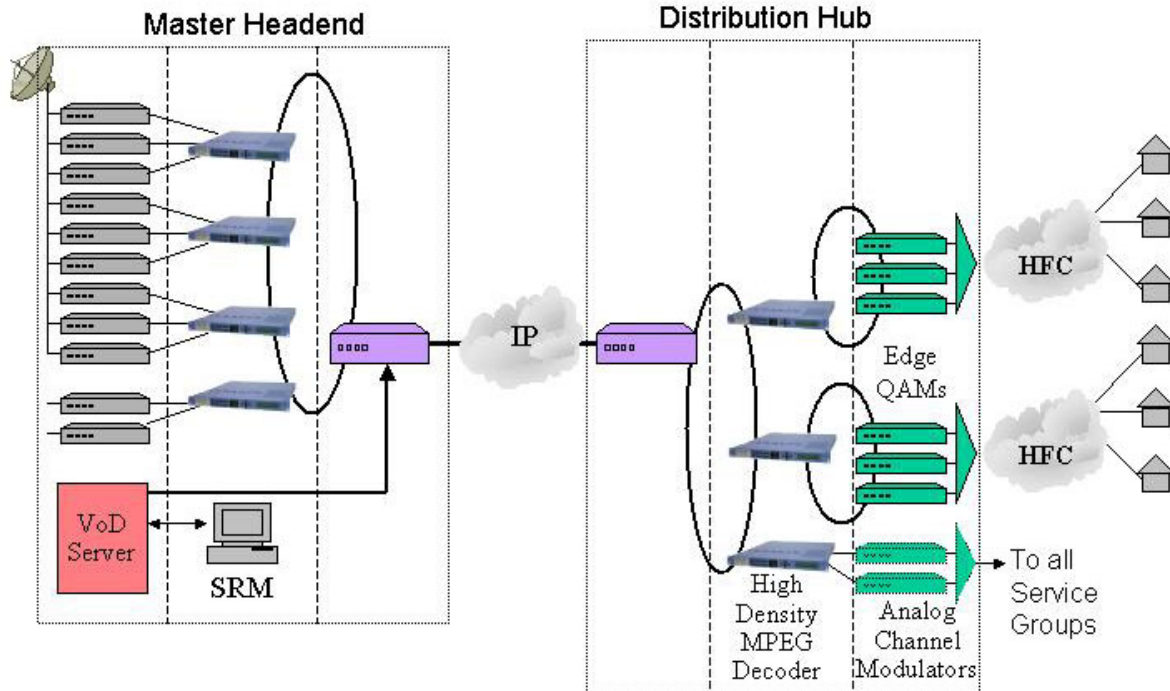


Figure 4: Digital Simulcast

While the operator still has to provide analog channels, an MPEG decoder and analog modulator are required for each channel to convert it back into NTSC format. Over a period of time, the goal of the operator will be to drastically reduce the number of analog channels because they consume so much HFC bandwidth compared with their digital equivalent. As analog channels are removed, more digital channels can be added for broadcast or on-demand viewing. Another advantage is that the DPI splicing in the digital domain is done before conversion back to analog, and so older analog insertion technology can be retired.

At this point in the migration path, we already have achieved complete convergence at the IP backbone level, with all four services – analog video, digital video, video-on-demand, and DOCSIS – running over a single backbone. This allows the operator to reduce operation costs and to operate larger networks more efficiently.

SUMMARY

This paper has illustrated the evolutionary migration path from today's hybrid transport model (using AM supertrunking) to a digital transport model that uses a converged IP backbone for all video, data, and voice transport.

REFERENCES

1. ANSI/SCTE 030 2001 (formerly DVS 380) Digital Program Insertion Splicing API
2. ANSI/SCTE 035 2004 Digital Program Insertion Cueing Message for Cable
3. ANSI/SCTE 67 2002 (formerly DVS 379) Applications Guidelines for SCTE 35 2001

(Available online from the SCTE web site at www.scte.org)

NEXT GENERATION VOD ASSET MANAGEMENT MANAGING THE VOD CHAIN END TO END

Chris Stasi - Vice President, Operations
TVN Entertainment

On Demand product has become a significant source of content to consumers but has only just begun to tap into its potential. As more and more content is accessed, stored and viewed in an on-demand environment, the requirements of asset management become exponentially greater. Systems will need to coordinate asset distribution from the time assets are exposed from a provider throughout their lifecycle, including distribution, ad insertion, streaming and financial reporting. One centralized system, managing content scheduling, offer creation, updates, distribution, tracking, metadata management, edge playout and wholesale reporting makes both MSO and provider VOD systems more efficient and scalable.

This paper will present an asset management strategy for both MSO's and providers covering the entire lifecycle of all pieces of the VOD architecture.

- *MSO asset tracking, updating and reporting*
- *MSO headend management*
- *Distribution management*
- *Content Provider Asset Scheduling, tracking, updating and remote monitoring*

On Demand content has only begun to scratch the surface of what is achievable in the consumer entertainment market. As the nascent technology becomes more and more refined a new level of interoperability and asset management is required. This includes not only the content itself but the associated

files and data as well. Currently the total hours of content available in the On Demand market is on the order of 3,000-4,000 per month. As the platforms become ubiquitous, and the technology easier for the average consumer to understand, the number of hours will increase quickly to 10,000 hours and beyond. Ominously, this does not scale linearly. With the introduction of the new elements and features of the Cablelabs standards and the rapidly growing incursion of advertisers into the On Demand market, the *type* of asset management required changes, as well as the amount.

Where We Are Now

Current asset management consists, predominantly, of two areas: metadata and content files. Both require a somewhat straightforward management style.

In the current Cablelabs¹ environment the metadata file is capable of a relatively small number of functions. It is the steward of the content from the time the content is created/encoded until EPG population. In the interim it guides the asset into residence on the VOD server and populates the associated databases. Subsequently its only functions are updates to certain allowable fields (on capable systems) and deletes.

The content file obviously functions solely to be viewed and the corresponding management of it is currently straightforward. It exists only as part of its own, single, offer/package and when the license window runs out it is deleted from

the video server. It is also tracked solely by references to it within its single metadata file. All of this is about to change and those changes will introduce an enormous layer of complexity into the On Demand environment.

Changes on the Horizon

The biggest issues facing On Demand players are Cablelabs upcoming 2.0² standard and the changing nature of the On Demand content itself. Both cause some of the same issues and many of those issues can be solved in similar ways.

The first asset management issue is the Cablelabs specification change. Whereas currently there is one offer associated with each content file, and each file exists in only its own, single, content package, in 2.0 the content file(s) each exist at their own level. There are one or more title/offer level metadata files tracking each asset and potentially multiple potential files tying each asset to other assets. Thus the single content asset that currently exists on its own with its own metadata file will now be able to sit on a server as part of many different offers and be acted upon in many different ways. This affects the content provider offering the content, the distribution company sending and tracking the file, the server allowing access to the file, the EPG's finding and displaying each offer of the content and the billing systems tracing usage. Many of the links in this particular value chain are not currently capable of the required changes.

The second issue is the changing type of On Demand content. The On Demand platform began mostly as a sister entity – programmatically – to the Pay Per View product. Movies were the first entry as well as Premium channel's SVOD offerings. Consumers were most familiar with this

product and had shown a willing appetite. Most of this programming was long-form and therefore small in number and easily trackable. Over the past 18 months the introduction of music videos, barkers, magazine shows and various other short-form assets has shown the direction content is headed. Advertising is the obvious next step and with its plethora of short-form, unique, content appears to be the 900-pound gorilla entering the room.

Macro Problems

So what asset management problems present themselves in this environment? There are a myriad of small ones, the devil is always in the details, but the major ones can be broken down into some sizable buckets:

1. The lack of a widely used, ubiquitous data set surrounding all On Demand content
2. Vendors within the On Demand space exhibiting different capabilities and requirements
3. Protecting the content to the satisfaction of content providers
4. The inability to update content in a near/real-time, meaningful way

Problem 1: Differing Data Sets

Potentially the biggest threat to widespread content usage is the inability for providers and MSO's to track content in a meaningful way. Currently the Cablelabs specification does a comprehensive job giving all parties the same ability to do so but not everyone is taking advantage of the information. The most widespread current Cablelabs spec is 1.1 and it has a fairly straightforward, well thought-out structure. Each content file, be it the main asset, a corresponding piece of artwork, or a trailer,

has its own asset ID. The metadata file corresponding to that asset and all the offer information it contains also has an asset ID, called Title ID. The combination of all this is the content Package and it has an ID of its own. The media files should keep their ID forever and any changes to the metadata should result in a change to the Title ID. Any change whatsoever to any aspect of the package results in a change to the Package ID. However this is not occurring across the industry. Widespread interpretive differences of the spec as well as embedded functional roadblocks have caused many providers as well as equipment and distribution companies to reach a different conclusion about how the data is managed. This has caused issues along the entire length of the On Demand chain. A provider has problems because if he/she does not track a content file with the same Mpeg ID there is no way to know how many times that same piece of content has been reintroduced or to get any meaningful comparative usage data back from an MSO. A distribution company cannot maintain any semblance of an organized library and distribution system if an ID structure is not followed for the life of an asset. A server vendor cannot reasonably be expected to cover all various types of content introduction, and make any type of meaningful error reporting on content, if it cannot expect a standard set of rules to be followed. If Titanic is introduced into VOD for the first time at one price, which filters throughout the server and billing interfaces, and the following year returns at a lower price, but the Title ID isn't changed, there is no standard way to report on the different price points. Data fields within the metadata files, such as Billing ID's, cause the same issues.

Solution:

So what can be done to make it easier to distribute and track content in an On Demand environment? The easiest and most obvious is for providers to all work off a common data system. Not that the onus should fall only on them but as they are the beginning of the road they get the first responsibility. The Cablelabs 1.1 format allows all providers as much flexibility as needed to identify their content. The use of domain name as the first half of a concatenated ID guarantees uniqueness across providers. The assignment of the ID's *within* each provider is then up to the provider itself to ensure and maintain uniqueness.

Following that, it is up to the distribution companies to maintain the provider's ID integrity throughout the industry. Maintaining this trackable ID structure is crucial to insuring that the content itself is traceable from a providers own management system, updatable from that same system and won't clash with assets being distributed to a location by another distributor.

At that location it is then up to the differing equipment vendors, mostly gateway, server and billing systems, to be able to ingest the ID string in the way it is intended. Currently there is a severe disconnect throughout the industry at this point and it causes not only a loss of traceability but literally a stoppage of content propagation and usage-data flow, which results in revenue loss. Following the protocols in the chart below shows the intended use of each ID level and what should and should not be passed.

	Original Offer	New Offer
Provider	TVN	TVN
Provider ID	TVN.com	TVN.com
Title	Spider Man 2	Spider Man 2
License Start Date	1/1/05	3/1/05
License End Date	1/31/05	3/31/05
Package Asset ID	TVNX1234500000000000	New
Title Asset ID	TVNX1234500000000001	New
Billing ID	56785	New
Movie Asset ID	TVNX1234500000000002	Same
Movie Content Value	TVNX1234500000000002 movie.mpg	Same
Preview Asset ID	TVNX1234500000000003	Same
Preview Content Value	TVNX1234500000000003 tr.mpg	Same
Poster Art Asset ID	TVNX1234500000000004	Same
Poster Art Content Value	TVNX1234500000000004 photo.bmp	Same

Upon the introduction of Cablelabs 2.0, the data set, if possible, becomes even more important. In 2.0 the elements exist on their own, apart from a singular, ever-present package ID. Collections come into play and they enable content assets to be acted upon by metadata that hasn't originated with the asset's initial introduction. A collection could consist solely of metadata introduced halfway through a movie's lifecycle offering a discount if the movie is viewed in tandem with a related new release feature new to the video server. If the collection metadata refers to an asset ID and that ID is not recognized the same way by every system nationwide the collection is not usable and again revenue is lost because of a disconnect on how best to interpret data.

The features and benefits of moving everyone into the same ubiquitous data set are easy to see. The most basic is simply the ability to track where your content is and who is using it. Building on that enables the introduction of ads. If you know where a hit movie resides and everyone knows the same ID structure it's easy to introduce an ad into a collection or playlist available to the end-user. It also enables the main content to be re-used without requiring re-pitching and

reintroduction to the server each time, saving money, time and processing power.

The widespread adaptation of the same data structure everywhere is a hurdle that has to be overcome quickly in order for On Demand content to evolve into the next generation product everyone desires.

Problem 2: Different Requirements from Edge Equipment

Similar to the previous problem, but potentially trickier to solve, is the differing level of capabilities and requirements of hardware throughout the On Demand architecture. These differences can be seen in areas as diverse as encoding specs, data requirements or EPG display capabilities, yet they all have a solution within the realm of asset management.

On Demand Asset management doesn't mean simply tracking data but also the construction and movement of the media itself. This area has encountered some difficulties because of different requirements across different platforms.

A perfect example is that the encoding specifications within Cablelabs³ have set the

baseline for asset construction but there are far more areas for missteps within the grey area of encoder setup itself. A configuration such as Program Stream has a choice within the specification but if an asset encoded on a valid, compliant stream arrives on some servers it will fail to ingest correctly even though it is a “legal” encode. Obviously a spec in a new environment cannot possibly be expected to have uncovered every nuance across all manufacturers’ equipment, but the introduction of elements such as Gig-E switches on the VOD network and other increasingly common architecture changes require a bit more standardization.

As mentioned in the previous section, billing systems currently present a few problems in the next generation On Demand environment. Many are set up to deal with ID’s from the PPV world. This was fine when there were assets numbering in the dozens arriving at a system each month and the only difference was time of day they were being watched. On Demand content has already reached the level of thousands of pieces of content per month and is growing rapidly. No provider is going to allow his or her content to be used forever without meaningful revenue in some form. That requires either direct usage-tracking for paid transactions or click-data for advertisements. Many current billing systems cannot provide this data based off the 1.1 and 2.0 ID structures. If this continues, the introduction of playlists, collections and increased shortform data will come to a standstill.

EPG’s also need to be able to offer the same minimum usage experience regardless of the platform they rest on. Currently each EPG is capable of different data capture and displays. While there will always be commercial reasons to offer different capabilities, a certain baseline set is

beginning to emerge through the Cablelabs group.

Solution:

So what can be done to standardize the requirements for different vendor equipment in the On Demand environment? Tough question. Anyone who has spent any time in the space knows that not only are the normal political elements involved in equipment purchase, vendor relationships and sales incentives but the element of speed to market also comes into play much more heavily in new technology centers. Manufacturers are forced to design on the fly and implement upgrades and rollouts before they’re ready in order not to lose market share to a competitor. But this shouldn’t stop the industry as a whole, through Cablelabs and other like organizations, from beginning to implement standards that go beyond the data and baseline encoding specs that are currently deployed.

While there will always be differences between rival commercial products, which obviously should be encouraged, there is no reason not to ensure certain quantifiable, measurable benchmarks be achieved in more areas. In order to be Cablelabs compliant a server and a distributor need to meet certain criteria in metadata creation and mpeg settings. Other settings need to be implemented as well, such as minimum capabilities on an EPG or minimum data compliance from the billing side. There is no reason a billing system cannot be certified as compliant in the On Demand environment the same way a server vendor must be. As more and more small MSO’s get into the On Demand space, and as more and more new companies bring their equipment into it, the natural propensity will be to become more disparate in the technology sets, not less.

Broadening the standards to include more equipment, encoding, display and interface requirements is the necessary next step to ensure advancement and interoperability.

Problem 3: Content Protection

As a content distributor, it has become obvious over the past few months that one of the biggest waves cresting on the horizon in 2005 is content protection and encryption. Providers are increasingly focused on ensuring their content is as safe as possible as it moves through its digital lifecycle. In an On Demand environment, where content is being passed from studios to distributors, sent over satellites across the country and taken down in hundreds of headends with almost as many differing security protocols, it's easy to see why a major motion picture studio or television network would be concerned about their content falling into the wrong hands. The seemingly endless downward spiral of profits within the music industry due to peer to peer file-sharing and illegal downloading is nothing if not alarming to content owners, and with good reason.

The question obviously isn't *whether* to protect and encrypt content but *how* to do it in a way that makes managing the protection protocols and the content itself as easy and as transferable as possible. Without some format for doing this, managing the asset throughout the VOD chain will not be possible.

The biggest problem implementing such a format is the number of places the content is touched in an On Demand environment in order to reach the end-user. It starts with the dub house, moves to one or more distribution/transmission companies where it is encoded and multicast and ends up at the MSO/Telco's headend. It is there that it

ingests into the video server and resides during its license window for streaming to the customer's set-top box. Every one of these touch points is susceptible to either the easy duplication/theft of content or at least the appearance thereof. The studios will not differentiate between the two and will require all touch points to have the same level of security.

The fact the studios haven't been more focused on this issue so far is due to a myriad of reasons, one of which is the seemingly safe environment of the secure headend. However this won't last. According to a study by AT&T and the University of Pennsylvania,⁴ 77% of trackable movie theft has been traced back to studio insiders. If the actual film companies themselves have been so compromised, an MSO headend, where the files are already digitized, is an even easier mark.

Solution:

So how do we prevent such a potential platform-destroying problem from taking root? The issue, obviously, is to protect the content as thoroughly and as long as possible. But with all the previously mentioned touch points for any single piece of content how does that take shape?

The best and most scalable solution is to encrypt any piece of content from its moment of creation (encoding) and offer the choice to leave that encryption on, or decrypt, at the natural handoff points throughout the distribution chain. With a modular, strippable solution in place the content is never clear-text except in a closed network, while it is simply changing cipher, not eliminating it. So how does this look?

There are multiple types of encryption and no two MSO's will want to implement it exactly the same way. Some will favor pre-encryption and some will favor session-based. Many factors will go into this, each as specific to the MSO, and as valid, as any other. The best way to enable all of them is to offer the content to them in a way that makes either possible.

A distribution company is generally the first player in On Demand to touch the content in a digital format. (This is a generalization as many content providers do their own encoding.) The first thing any distribution company must do is create an MPAA approved library storage system that is approved from a physical access standpoint and also stores encoded files solely on a private network. This obvious step will keep access to the files to as few people as possible.

At this point the files should be encrypted, but how to do that this early and still fit into the modular scheme mentioned above? A standard 128 bit AES encryption scheme can be applied at this, or any, point and passed down the chain. As the content is received at an MSO's headend the receiving device can be configured to either decrypt the content prior to handing to the video server or pass it through in its encrypted state. If the decryption takes place the content must be handed to the video server across a private network to ensure its safety. If the MSO chooses they can receive the content still in an encrypted state and pass it all the way to the set-top. This of course requires the set-tops to host a decrypt client that can receive the decryption keys. It is in the MSO's best interest to get set-top manufacturers to certify as many of these clients as possible to ensure competition and an open standards-based solution.

The encrypted data packets are transparent, so firewalls, proxies etc. can see the encrypted packets and pass them through as if they were unencrypted. This allows any encrypted file to pass itself down through the value chain and be acted upon (transferred, streamed, trick-file creation) as if it were clear text. If this end-to-end pre-encryption is carried through it is up to the distribution company to manage the key-list server and enable all of its MSO's to receive these keys. It also has to ensure that the encryption scheme is standards based. Getting all the encryption companies to maintain a standard scheme, and potentially a universal cipher, is underway and will go a long way to making the transition to a completely encrypted product easier.

If an MSO chooses to receive the content as clear text the obvious choices are session-based encryption, with its higher degree of safety but commensurately higher costs, or to re-encrypt on the server. This requires the same client decryption on the set-top.

There are many complications as encryption enters the On Demand environment. By adhering to a standards-based, modular, flexible encryption scheme that allows the cipher to be removed or passed through at any point, the studios and networks will be satisfied their content will be protected as thoroughly as they demand. This will remove a major roadblock from managing On Demand assets throughout their lifecycle

Problem 4: Immediate Data Access

A problem that has not derailed the rollout of On Demand content but will soon pose a big problem is the inability for programmers, distribution companies and even MSO's to act upon the data surrounding their content in a real-time,

meaningful fashion. As any new technology rolls out to customers, the wow factor carries it through its initial hiccups. This fades quickly and soon the product needs to meet consumers' up to the minute demands or it will fade into obscurity. Once the technology is settled the marketing takes over. The On Demand product chain has been very limited in what the marketers can do because the data they can get their hands on has been very delayed. Even the small amount that comes to them quickly can't be updated or changed in an easy, ubiquitous fashion. This will have to change quickly for the On Demand product to gain relevance to the end-user. It will also be a requirement for advertisers. Nobody wants an ad running that is out of date by the time he/she has the ability to update it.

With the upcoming release of the Cablelabs 2.0 standard the ability to wrap metadata around content differently will help bring some marketing capabilities into the On Demand environment. An advertiser can use either a collection or a playlist to attach ads to free On Demand content or aim targeted ads at a willing viewer. An asset price or window can be changed, or that asset can be packaged with like material to use one as a loss leader, after the material has been in its viewing window for some time. There are still some major hurdles that need to be overcome before this is anything more than a technical capability though. There is no performance data available to providers to determine how, for example, a New Release movie is performing from market to market until the early heavy-usage period has passed. Even with this data at their fingertips the providers have no way of acting upon it remotely without involving another vendor or distribution partner. The only way to aim these ads correctly, or know what offer to adjust on a new title, is to have immediate access to usage data and a system

to take advantage of it.

Solution:

The easiest was to prevent this data restriction from adversely affecting the programming in On Demand is to open up the usage data and give everyone along the distribution chain access to changing it. This consists of two major steps.

The first is to allow a standardized interface into the server usage data. This doesn't necessarily mean giving providers or distributors private customer information. What it does mean is at least passing back generalized buy rates and click data. This can be done through a standardized reverse interface from the server outward. A distributor can link into the server and extract (or more likely be fed) a set of usage data that can filter back to the distributor for storage in their database. This data can then be mined to determine what pieces of content are performing better or worse than expected. Then, assuming the MSO's equipment is updatable (which by CL 2.0 it would have to be) the data can be used to extend, change or refine the offering to the consumer, making the asset more valuable. This interface has begun to get some traction within the Cablelabs consortium.

This information will be worthless however without the ability to view it, make decisions and act upon it. A distributor receiving this information from a server needs the ability to grant any provider using his/her system access to the data and update functionality in real time. A distributor's system needs to be able to author content within it and send that content to the destination headend's but that's not all. In order to make data management meaningful, any data brought back to the provider

through the reverse interface needs to be editable by the same provider without having to go headend by headend. One centralized system, able to update data, nationally or singularly, is necessary. This system needs to manage content preparation, transmission and headend interfacing at every level, in other words an overall layer sitting on top of the entire value chain. This enables any interested party, anywhere on the value chain, access (with permissions of course) to their content, be they provider, distributor or MSO. Without this universal overlay the job of manipulating content will grow exponentially larger and quickly get out of hand. This will have the effect of making content quickly grow stale, thereby seeming less interesting to the consumer. As advertising and short-form assets continue to be a larger part of On Demand content and as the number of files associated with each piece of content grow, one system capable of managing this content from end to end is a necessity.

Moving Forward

On Demand (in one form or another) will eventually be how most people watch most content. Many more issues than discussed here will eventually pop up, some of which take everyone by surprise. But planning for the next generation of asset management starts now. Done right, everyone involved can benefit, from consumers through the smallest distributor.

REFERENCES:

1. <http://www.cablelabs.org/projects/meta-data/specifications/specifications11.html>
2. <http://www.cablelabs.org/projects/meta-data/specifications/specifications20.html>
3. http://www.cablelabs.org/projects/meta-data/specifications/content_encoding.html
4. http://www.aeanet.org/GovernmentAffairs/gamb972_ATTReport_MoviePiracy.asp?bhcp=1

OPTIMAL QAM ASSIGNMENT IN THE PRESENCE OF MIXED SD AND HD STREAMS

Jiong Gong and Yasser Syed
Cable Television Laboratories, Inc.

Abstract

When VoD systems need to accommodate both SD and HD streams, the traditional capacity engineering rule for deploying QAM modulators face a new challenge. Two issues arise from this new paradigm. One is the streaming bit rate for HD, and perhaps both HD and SD streams to optimize the system performance. The second issue is the QAM allocation algorithm to minimize system inefficiencies. We propose a solution that has the potential to significantly improve system performance to accommodate a mixture of SD and HD VoD streams compared to prevalent methods.

INTRODUCTION

With the rapid penetration of HD TV sets in the consumer electronics market, cable companies have been aggressively deploying High Definition (HD) cable channels in recent years. A powerful competitive response to the DBS offerings is the HD VoD service. When the VoD streams consist of only one type that is of Standard Definition (SD) TV, the Quadrature Amplitude Modulation (QAM) assignment algorithm is not critical in affecting overall system performance. Commonly implemented algorithms include starting from the busiest QAM and starting from the least-busy QAM. Either way is not going to affect the blocking rate of the system. To make full use of a QAM channel capacity, bit rate is often chosen so that all possible numbers of streams in a channel form a congruence class

of a modulo in the streaming bit rate with the division remainder to be zero.

When VoD systems need to accommodate both SD and HD streams, the traditional stream capacity assignment rule for QAM modulators faces a new challenge, in that it is possible to incur blockage on each of the QAM modulators while jointly they have the capacity to support the arrival of a new stream request. In other words, the posterior allocation of the streams is suboptimal. Both the busiest and the least-busy rule tend to have suboptimal allocations. This brings about the issue of finding an alternative algorithm to improve the allocation efficiency. We propose in this paper an optimal solution that is a significant improvement over current methods.

VOD STREAMING BIT RATES

Before getting into QAM allocation algorithms, it is important to first look at the issue of VOD streaming bit rates. Design of bit rates should consider both the issue of fully using QAM resources and the quality issue perceived by viewers. For each type of stream to fully utilize the useable bandwidth capacity of a QAM, there must be a modular relationship between each type of stream. Additionally, if a different capacity QAM is used then a modular relationship must still also exist for full utilization. In reality, the limits of this relationship are dependent on the modulation types in use. Since quality is dependent on bit rate, there is an additional limit to designing the modularity factor between the two types of streams.

There are some real practical limitations on this relationship. The data rate used in a QAM-designated VOD service today is 37.5 Mbps for 256 QAM and 26.25 Mbps for 64 QAM. The additional capacity in the data rate is reserved for in-band traffic. Today a typical SD VOD stream is at a constant bit rate of 3.75 Mbps which is good quality for MPEG-2 paid movie content. If one wanted to determine the HD MPEG-2 stream bit rate based on modularity on a fully-utilized QAM, then HD data rates would be 7.5 Mbps (2x of SD stream) or 18.75 Mbps (5x of SD stream), which is either too low in quality or too much data rate used. Alternatives to this would be 11.25 Mbps (3x of SD stream) or 15 Mbps (4x of SD stream). Both of these offer acceptable quality, but they are not fully-useable “pure” QAMs, requiring either a 2 HD/2 SD- 256 QAM (or 1 HD/2 SD- 64 QAM) stream configuration or 3 HD stream/1 SD stream-256 QAM (or 2HD/1SD- 64 QAM) configuration because of a modular relationship.

In typical systems, a VOD service to a node is allocated in 4 QAMs (for now let's assume 256 QAM) or an integer multiple of it. This has to do basically with fiber distribution to a node. For a pure SD VOD service, this would be about 40 streams that could service about 400 homes, assuming a 10% peak capacity. For a pure HD VOD service, this could be either 8 to 12 streams that could service from 80 to 120 homes, assuming a 10 % peak capacity. If a QAM allocation algorithm is not properly configured, this would either lead to blocking of HD streams to 4 to 8 streams by the wrong placement of 4 SD streams. That leads to only 40 to 80 households for HD VoD service, assuing 10% peak usage rate.

The amount of time of this blocking would depend on length of overlap that the bandwidth is reserved for each movie. The

current pratice in many VoD systems is to tear down the stream, if the stream incurs more than 5 minutes of pausing. One futher complexity is whether the torn down stream needs to have any priority over new stream requests, if it needs to be resumed again. If so, this would affect how the bandwidth is allocated and the amount of time the bandwidth is reserved (e.g. a 2 hour movie may typically have a reserve time that has an extra 20-30 minutes).

QAM RESOURCE ALLOCATION

In this section, we describe a QAM allocation algorithm that we believe represents a significant improvement over current prevalent methods. We start off by describing a mathematical framework to model the problem. Suppose a collection of n QAM modulators is deployed to serve a VoD service group. Let q_i , $i = 1, 2, \dots, n$, denote the used capacity of each QAM modulator i . Total capacity, Q , is assumed to be the same for all QAM modulators. Therefore, the remaining capacity that can be used for new stream requests on that QAM modulator is then $Q - q_i$. Let r_s and r_h denote the streaming bit rate respectively for SD and HD streams. The two types of streams may arrive at a collection of QAM resources according to two distinct random processes, such as the Poisson process, but exit the system based on the same holding time distribution. We call the snap shot state of all QAMs $[q_i]$ at a particular time as an allocation. We define an allocation as inefficient, if,

$$Q - q_i < r_h, \forall i, \text{ and } \sum Q - q_i \geq r_h$$

In other words, none of the QAMs individually has the capacity, even though the sum of all available resources on each QAM is able to support a HD stream request.

Note that while each type of stream is in itself modulus in its own bit rate, they jointly are not when they are mixed together in a QAM. As a result, inefficiency tends to arise when streams are mixed together. Both busiest and least-busy algorithms tend to create mixing QAMs (i.e., both SD and HD streams carried by the same QAM) as the joint process of the two stream types is a mixture of two random processes. Naturally our improvement over the current methods is in the direction of minimizing the probability of a mixing QAM.

The state of each QAM modulator can be categorized into four possible types:

- It is entirely empty.
- A mixture of SD and HD streams are occupying it.
- Only SD streams are occupying it.
- Only HD streams are occupying it.

Mathematically we denote these four types accordingly by defining a state function as:

$$S_i(q_i) = \begin{cases} 1, & \text{if } q_i = 0 \\ 2, & \text{if } q_i = x_i r_s + y_i r_h, x_i \neq 0, y_i \neq 0 \\ 3, & \text{if } q_i = x_i r_h, x_i \neq 0 \\ 4, & \text{if } q_i = y_i r_h, y_i \neq 0 \end{cases}$$

where x_i and y_i are positive integers representing the number of SD and HD streams occupying QAM modulator i . In the above four states, we call a QAM in state 1 an empty QAM. We call a QAM in state 2, that is $S_i(q_i) = 2$, a mixing QAM. QAMs in state 3 and 4 are called non-mixing SD and HD QAMs respectively.

Since our algorithm is based on the principle of minimizing the probability of mixing two types of streams within a QAM modulator, it is then straightforward to

prioritize over the states of available QAM modulators for stream allocation. The first priority is to go with a non-mixing QAM of the same stream type. The next priority is to go with an empty QAM. Then go with a mixing QAM. The last resort is to create another mixing QAM – going with a non-mixing QAM of the opposite type.

If there are multiple QAM modulators available within the same state class, priority is given to those QAM modulators that have a larger likelihood of becoming a non-mixing QAM or an empty QAM once some streams start to drop. This implies the following rules:

If multiple non-mixing QAMs are available to a stream request of the same stream type, priority should be given to the busiest non-mixing QAM, because other non-mixing QAMs have a higher likelihood of being empty.

If multiple mixing QAMs are available to a SD or HD stream request, priority is given to the busiest mixing QAM, because other mixing QAMs have a higher likelihood of being non-mixing or empty.

If multiple non-mixing QAMs are available to a stream request of a different type, that is if a stream request will have to create a new mixing QAM, priority is given to the least busy QAM, because it has the highest likelihood of becoming non-mixing again.

With these rules explained, the selection algorithm of a particular QAM modulator by a new stream request is then based on the following sequence. We take a SD stream request as an example.

1. Identify a set of I , s.t. $Q - q_i \geq r_s$ for $\forall i, i \in I$

1.1 If I is empty, reject the stream request;

2. Identify a subset of J , $J \subseteq I$, s.t. $S_j(q_j) = 3, j \in J$;

2.1 If J is empty, go to the next step;

2.2 If J has multiple elements, select

$$j^* = \arg \min_{j \in J} Q - q_j;$$

2.3 If there are multiple j^* , select randomly among j^* ;

3. Identify a subset of J , $J \subseteq I$, s.t. $S_j(q_j) = 1, j \in J$;

3.1 If J is empty, go to the next step;

3.2 If J has multiple elements, select j^* randomly;

4. Identify a subset of J , $J \subseteq I$, s.t. $S_j(q_j) = 2, j \in J$;

4.1 If J is empty, go to the next step;

4.2 If J has multiple elements, select

$$j^* = \arg \min_{j \in J} Q - q_j$$

4.3 If there are multiple j^* , select randomly among j^* ;

5. Identify a subset of J , $J \subseteq I$, s.t. $S_j(q_j) = 4, j \in J$;

5.1 If J has multiple elements, select

$$j^* = \arg \max_{j \in J} Q - q_j;$$

5.2 If there are multiple j^* , select randomly among j^* ;

We conclude this paper by presenting the major result of the paper as in the following proposition. The result waits to be verified by simulation tests. The test scheme will generate two arrival random processes and document blocking rates for the three

alternative algorithms. If our algorithm generates least blocking under the same load, it would then verify the theoretical result.

Proposition: The QAM allocation algorithm described above is more efficient than the busiest and the least-busy algorithms.

Proof: Suppose the system has only empty and non-mixing QAMs. At this point, the system has capacity to accommodate new streams. As more streams are added to the system, the algorithm can only incur mixing when the last QAM is called for to meet the demand. In other words, mixing occurs when the system is close to full capacity except the last QAM. Note at this point, all other QAMs are non-mixing. Let j denote the only mixing QAM. When it is full, $Q - q_j < r_h$ implies $\sum Q - q_i = Q - q_j < r_h$. This is because all other QAMs are non-mixing, and each type of stream is modulus in its own bit rate. Therefore there would be no inefficiency according to our definition. Therefore our algorithm is more efficient.

On the other hand, mixing of QAM in the busiest and the least-busy algorithms is a random event. Therefore inefficiency is likely to result with a higher probability.

CONCLUSIONS

This paper presents an alternative QAM resource allocation algorithm to accommodate a mixture of SD and HD VOD streams. Our analysis indicates the need to first design the stream bit rates so as to be modular with respect to the full QAM capacity, as well as to be modular with respect to each other. The subsequent allocation algorithm then calls for avoiding the mixing of different types of streams to the extent possible. This principle is very much

like the principle of ocean freight container shipment, where modularity is critical in making full and efficient use of the ship capacity.

In the future when there are more types of streams in the same QAM set (e.g., alternate codec versions like MPEG 4-AVC or VC-1, and each with an SD or HD version), one should consider the modularity for selecting the bit rates for these. Additionally if we consider supporting VBR streams instead of CBR streams. Some parameters may need to reflect modularity, start time difference in QAM streams, and value of the streams.

PREPARING FOR THE NEXT GENERATION OF VOD TECHNOLOGY – A CONTENT PROVIDER 'S PERSPECTIVE

Dave Bartolone, Vice President, Technology
TVN Entertainment

Abstract

Video on Demand has advanced from a simple process of managing 50 hours of movie content with 7 day lead times to a robust collection of long form and short form content with lead times reduced to hours instead of days. Add to this the next generation of VOD functionality that will provide capabilities that could increase the amount of assets to be distributed and managed significantly putting greater demands on providers' Asset Management and Distribution Systems and processes.

This paper explains and illustrates how VOD has changed in the past 4 years and where it is headed based on certain trends and new capabilities introduced by the next version of the CableLabs VOD specification. This paper examines the impact these advances will impose on content providers and how they can prepare for it in the future.

BACKGROUND

In the past 4 years, TVN has seen the amount of VOD content distributed grow at a steady and predictable pace. In April 2001, TVN began pitching VOD content to its first commercial VOD site with approximately 50 hours of movie content with an average run time of around 100 minutes per title. That was soon followed by the launch of approximately 300 hours of SVOD content from the premium networks with average run time of around 70 minutes

per title. Then, in October 2002, TVN began including VOD content from cable network providers. We ended the year with approximately 250 hours of free On-Demand content with average run times of approximately 54 minutes each. Today, the TVN network distributes over 3000 hours of content a month from these and other VOD categories with average run times of approximately 49 minutes each.

Delivery lead time is defined as the amount of time that a title must be pitched from TVN and caught by a VOD site before the view start date of that title. In the beginning, having only feature and library content to deliver, the VOD delivery process required delivery lead times ranged from 7 to 10 days. As more time sensitive content, such as current events, highlights from recent sporting events, etc., was added, we now have some delivery lead times reduced to as little as 6 to 12 hours.

Most of TVN's content growth is a result of adding over 70 content providers to the mix of aggregated programming available through TVN. Under the current CableLabs 1.1 VOD asset structure, where a package contains one metadata file, one to two MPEG files and possibly a graphics file, this volume equates to between 5,000 and 10,000 files a month to be distributed and managed by TVN. If you consider that at any one time, these 5,000 to 10,000 files may be multicast to between 80 and 100 sites at once, that equates to an average of 675,000 distributed files managed per month.

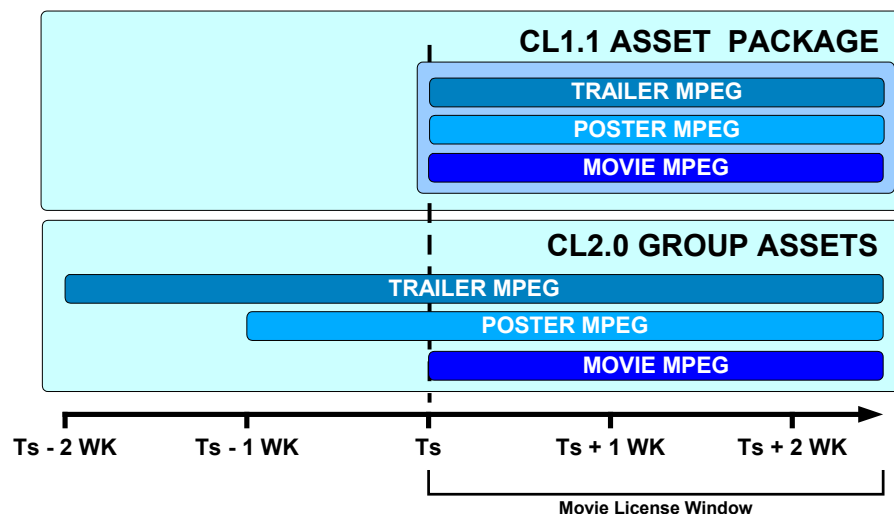
Knowing the current simplicity of a VOD offering that contains a single version of a title at a single price, we have been able to accurately predict the rate of increase and scale up our support systems and operations accordingly. Our next task, however, is to examine the new capabilities introduced by the next generation of VOD and plan accordingly.

NEW CAPABILITIES FOR VOD

Currently, the majority of VOD installations operate under the bounds of the CableLabs 1.1 VOD specification (CL1.1). The next version of the specification, CableLabs 2.0 (CL2.0), will introduce much needed flexibility in forming, distributing and offering content. Following are a few examples:

Breaking up Asset packages - An intended feature of the new specification is to allow a content provider to break up a CL1.1 asset package into individual elements and assign individual window dates to each element. Each package element can now be introduced to a VOD system independently and live on its own. For instance, a content provider would have the capability to efficiently distribute and display the preview or poster art of an offering before the movie is made available for viewing.

The diagram below shows the difference between the capabilities using CL1.1 versus those capabilities using CL2.0.



Under the CL1.1 scenario, the viewable window dates of each element of the asset package opens at the same time at T_s . However, under the CL2.0 scenario, in this example, the preview and its metadata is introduced at $T_s - 2$ weeks, the poster art and its metadata is introduced at $T_s - 1$ week, and the movie asset is introduced at T_s .

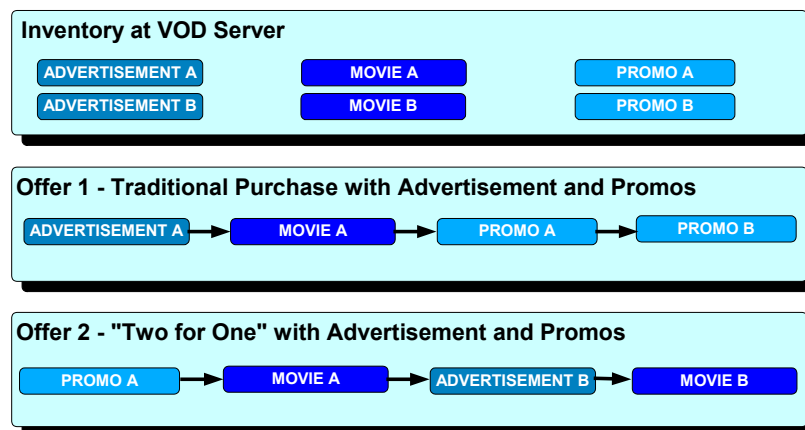
This simple but useful concept will allow for the promotion of upcoming VOD titles with the goal of increasing awareness and buy rates for the movie asset.

Using Playlists to Insert Ads - The playlist concept has been around for many years and is a very versatile tool that will allow content providers, and eventually consumers, to specify a list of individual MPEG files and a sequence in which they

should be played out. This one feature alone will create an enormous opportunity to dramatically change the way VOD programming is being offered today. Even though there is more work to be done to define business rules, resolve potential billing system issues, etc., playlists offer a huge potential in the ability to create unique and appealing offers to consumers. One opportunity for both content providers and MSOs is the ability to gain additional revenue through placing advertising within a VOD offering. Additionally, a single asset can be shared across many playlists allowing for create marketing and pricing discount schemes such as “two for ones” and “many for one” offerings based on a common

attribute of each piece of content such as genre or actor. In fact, the person defining a playlist can pick from a number of assets that have been placed on the VOD server, and, within the bounds of business rules agreed to by the Content provider and the MSO, create a multitude of different offers.

For example, you might have two ads, two movies, and two promos that a content provider has placed onto the VOD server. In Offer 1, you can combine one of the movies and two promos to make an offering. Or, in Offer 2, you can combine a different ad with a different promo and two movies to create a “two for one” offering.



This is an example of a static playlist that has been pre-determined prior to the view start date assigned to the playlist. Future versions of this feature may include the ability to create dynamic playlists based on subscriber viewing behaviors or stored demographic information.

Other uses for playlists include the formation of both category and main page barker videos. In this scenario, instead of editing a series of movie trailers into one video that then gets encoded, a playlist enables content providers and MSOs to reuse existing movie trailers to be played back, as a barker video, in either a set or

random order. Furthermore, since a playlist will be an asset having their own window dates, different barker videos can be applied at different times of the day or week.

Adding Chaptering Information – Similar to a DVD presentation, an intended feature of the new specification is to allow for each asset to be indexed and then have chapter information and chapter graphics applied to the indexed locations. Once a subscriber has navigated to and ordered a particular title, this feature will allow the viewer to now navigate within the asset. Viewers would also be able to skip past or skip to chapters as desired.

Managing Menu Categories – Currently, content is displayed in menu categories with no regard to the display order other than alphabetically by title. An intended feature of the new specification is to allow content providers and/or MSOs to define a list order of the menu category contents as well as time durations for the listing within the menu category.

Adding Keywords – In an environment where the amount of content on a VOD system exceeds the practical navigational ability of the subscriber, keyword search functions will become an attractive feature. In this scenario, a user can enter keywords, perform a keyword search, and then select from a list of results that might not have otherwise been displayed on the VOD user interface.

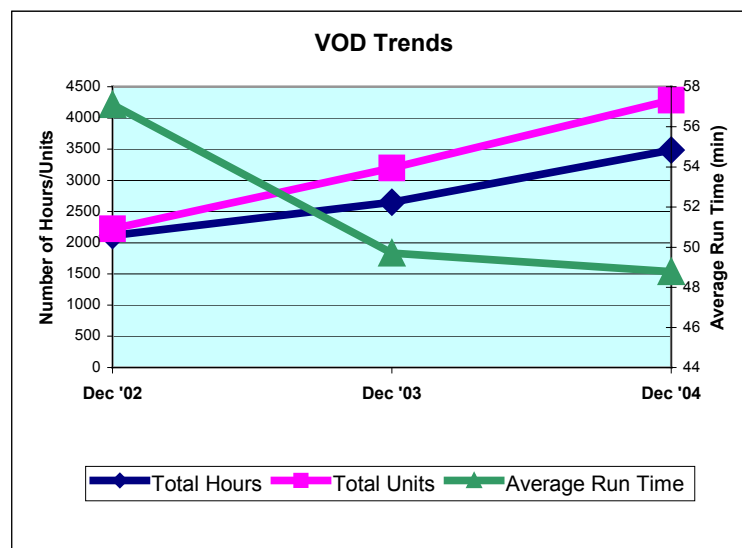
FUTURE TRENDS

New Requirements based on Current Capabilities

When considering the challenges of Asset Management and Distribution, content providers must consider the number of files, or units, distributed rather than the number

of hours. In other words, it may take four times the effort to distribute and manage four 15 minute titles than it does to distribute and manage a single one 1 hour title. By studying the trend of VOD operations using the current CL1.1 specification over the past three years, we found that the average run time per title has been steadily decreasing which inversely increases the number of units per hour of content to be distributed and managed.

The following chart shows TVN data gathered over the last three years illustrating how the average run time per asset has been decreasing. In December of 2002, TVN managed 2100 hours of in-window assets with an average run time of 57.11 minutes resulting in approximately 2225 titles. In December of 2003, TVN managed approximately 2650 hours of in-window assets with an average run time of 49.71 minutes resulting in approximately 3200 titles. In December of 2004, TVN managed approximately 3500 hours of in-window assets with an average run time of 48.78 minutes resulting in approximately 4200 titles. This equates to a 17% decrease in average run time over the three year period.



Another trend that we have noticed over the past year is a decrease in the delivery lead times. This is due to the introduction of time sensitive content such as sports highlights and current events.

If these trends continues, which we think they will, and there are no new capabilities introduced, then an unprepared Asset Management System will be challenged to distribute and manage this increased unit volume of content quickly and efficiently. Adding on top of that the new intended capabilities introduced by CL2.0, you may find an exponential increase in the number of units per hour to manage.

New Requirements Based on New Capabilities

Following is a review of some of the increased capabilities introduced by CL2.0 from an Asset Management and Distribution perspective:

Breaking up Asset Packages – Using the example of a CL1.1 asset package containing a movie file, a preview file, and a metadata file, the management of the asset package consists of the management of three individual files. Breaking up the package into multiple distributed elements of the preview one week, the art work two weeks later, and the movie asset as the final delivered asset will result in the creation, distribution, and management of six individual files. The result is a two fold increase in the number of units to distribute and manage given the same run time.

Using Playlists to Insert Ads – Under the current proposed specification for playlists, there is no technical limitation to the complexity and resulting amount of assets that can be strung together to form a playlist. However, there will no doubt be practical

limitations introduced during implementation. If we look at the most basic, easiest to implement playlist, it would consist of one or two promos together with a main movie asset. In this scenario, there are no dependencies on back-end process such as those associated with ad-viewing confirmation and reconciliation. Using this example, given a slight increase in total run time, up to two additional metadata files and two MPEG files will need to be distributed and managed. The result, again, is at least two fold increase in the number of units to distribute and manage given a similar run time.

Adding Chaptering Information – Although the indexing tags for the location of each chapter will be contained within the existing metadata file, the addition of optional chapter graphics will increase the number of units to be managed in direct proportion.

MANAGING ASSETS

New Demands for Asset Management Systems

It may appear obvious that the new capabilities introduced by the new CL2.0 specification also introduces new demands for Asset Management and Distribution systems. With asset packages being broken apart, tracking the relationships of these separated assets will be a key feature of any future Asset Management and Distribution system. New systems must keep track of each package element each having their own window dates that must all overlap at some point in order to produce a viable offering to a subscriber. Furthermore, as individual elements of playlists are identified, their viewable window dates must, at minimum, match the start and end view dates as defined in the playlist.

What Content Providers can do to Prepare

Although the timeframe for implementing these new features is uncertain, it is important for content providers to prepare for them in advance. Following are a few suggestions:

- Asset Management - Approach the new capabilities from an Asset Management perspective paying close attention to the challenge of managing more units per hour and what may seem to be disparate assets that come together on a VOD server through the use of playlists.
- Segmenting Files - For the short term, and only after playlists are implemented, consider separating content into segments so that other elements can be inserted in between.

- Adding Keywords – Start creating and adding keywords to your content.
- Chapter Graphics – Start identifying and storing graphics that can be used to identify chapter locations.

CONCLUSION

The new capabilities brought on by the introduction of CL2.0 are exciting and are bound to launch the VOD product into its next generation of usefulness. However, one must not underestimate the demands the new capabilities will have on Asset Management and Distribution systems. The sooner content providers can intelligently anticipate the use of CL2.0, the better they can prepare by upgrading or selecting an Asset Management and Distribution system or service that is capable of satisfying these new demands.

RF FINGERPRINTING: AN OPERATIONALLY EFFECTIVE METHOD TO REDUCE CABLE TELEVISION SIGNAL AND EQUIPMENT THEFT

Lee Pedlow
Sony Electronics, Inc.

Abstract

Presented is a robust method to self-detect the unauthorized relocation of digital cable television appliances, especially one-way or CableCARD-based devices, as a deterrent to signal (service) and MSO-provided equipment theft. The method offered has high resolution, yet requires no additional hardware to be added to the products in which it is implemented. Implementation of the concept uses an innovative application of resources already present in all digital cable compatible devices plus the real-time analysis of characteristic data obtained by the subscriber device through direct observation of its environment.

A system for automated device management is also presented wherein subscribers could self-activate attached devices without need for manual intervention by the cable system operator under normal circumstances.

INTRODUCTION

Digital Cable Appliances

Digital cable television appliances are becoming mainstream devices in the modern home. These devices may be stand-alone “set top boxes” that are either leased from the cable operator or purchased by the consumer through retail channels. Alternatively, the functionality of a digital cable appliance may be integrated directly into new television receivers as part of the “plug and play”

initiative for digital television, as mandated by the FCC.

As the cost of implementing digital decoding capabilities in consumer products rapidly declines and the prevalence of digital programming on cable television systems grows, the cable industry is marching toward removal of all remaining analog television services from their systems to reclaim spectrum, reduce operational costs and reduce signal theft.

In the foreseeable future cable operators will need to supply their existing customers having legacy analog televisions, VCRs, etc. large quantities of digital converters in 3to maintain continued operation of those analog devices in an all-digital network and as part of the operator’s compliance with federal regulations.

Industry estimates indicate that there may be four or more legacy analog devices attached to the cable system in a typical household in addition to existing digital cable converters for premium service access and CableCARD enabled products. Because of the sheer volume of digital converters that the cable operators will need to deploy in support of the analog devices presently in their subscribers’ homes and the fact that analog devices connected directly to the cable network do not pay for advanced services such as electronic program guides, video-on-demand or pay-per-view services, cable operators have no method to recover the huge additional capital outlay required

for supplying the advanced, two-way digital cable boxes currently available.

As a result, attention is now focusing upon providing very inexpensive, one-way digital converters for this purpose, delivering current analog subscribers like-for-like digital service at a significantly lower cost to the operator than would be encountered using the presently available advanced two-way devices. The cable operators, for regulatory and other reasons, intend to provide these one-way converters at no additional cost to their subscribers and believe that the cost of providing these devices can be more than offset through the recovery of valuable cable spectrum, elimination of signal theft and reduction of operational costs such as truck rolls for service connect/disconnect.

These simple digital converters are intended only for the most basic of service tiers, ones that are presently delivered in analog form and therefore have been left unprotected against unauthorized reception, unlike the current premium services which employ modern digital encryption. A conservative estimate is that more than one third of the channels carried in modern cable systems are presently analog basic services.

A May, 2004 press release from a major cable system operator, Cox Communications, indicates that roughly 11.5 million U.S. households steal these cable services each year at an industry cost of \$6.5 billion in lost revenue annually.

The transition of the basic subscription tiers from analog services to exclusively digital services having encryption applied will eliminate most of the present forms of signal theft that occur because these new digital converters will be individually addressable by the cable operator. Unlike today, merely having physical access to the

cable signal either through an unauthorized connection by tampering or because there hasn't been a costly dispatch of the cable operator's field personnel to the premises to implement a disconnect will no longer suffice for basic tier customers to receive services for which the cable operator is not compensated. This also applies to new digital television receivers if the owner has not obtained a CableCARD from the cable operator and had it electronically authorized for service.

The typical conversion scenario for the all-digital transition would be for a cable operator to upgrade a headend serving a community or city to carry basic tier content in digital form in addition to the analog format presently carried. Next, all current two-way devices deployed for decoding premium digital services are provisioned with new channel maps, directing them to receive only digital forms of content, including the new digital replacements for the analog tier instead of the present mixed format. In parallel, the operator will begin distribution of the new, low-cost one-way converters to existing subscribers based upon the number of cable outlets in the home that are reported by the subscriber as connected to a legacy analog device (VCR, TV, etc.). There is no way for the cable operator to determine the analog device count in a home without either surveying the subscriber or performing a physical audit inside their premises. The operator will likely deploy these new converter devices en-masse as each node served by a cable headend is converted from mixed analog/digital format to all-digital through the removal of analog services. A network node typically serves from 500 to 2000 customers and the converters must be available to all subscribers in a targeted node prior to cutover in order to avoid service interruption

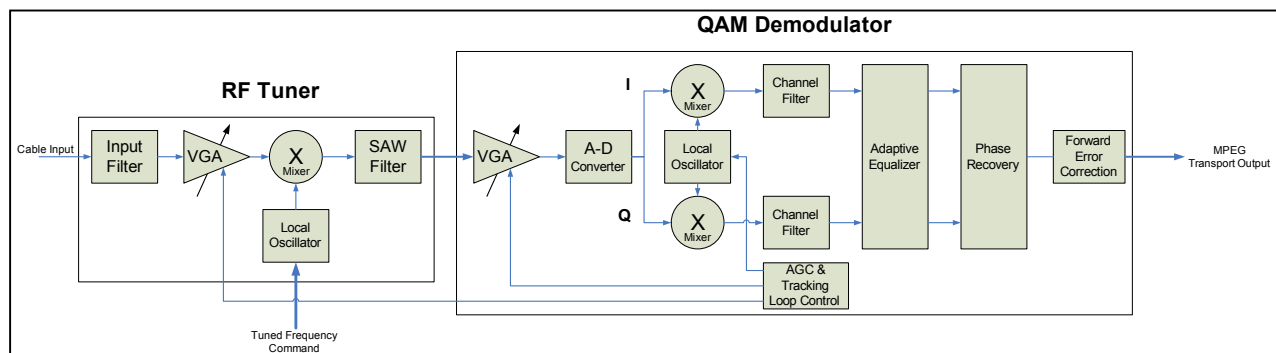


Figure 1. Generic digital cable network interface

While the introduction of all-digital services and low cost digital converters would seem to address all the issues of unauthorized viewing and signal theft, a new opportunity to deprive cable operators of fair payment for service emerges, made even more challenging because these low cost converters are likely only one-way devices,. When subscribers are contacted to determine the quantity of converters necessary for supporting the analog appliances in their home, the subscriber may intentionally “over-report” the quantity of analog appliances in the home. They can later provide the excess converters received from the cable operator to friends, family, etc. to “split the costs” of basic cable service. This is but one example of how one-way converters can be redistributed without the knowledge or consent of the cable operator in order to deny the operator of payment.

Since these new devices are issued by the operator to a valid subscriber, they remain authorized, but the cable operator is deprived of the full value of subscription revenue because the devices are present in locations other than the home of record for the authorized subscriber and the operator is therefore not compensated by the additional, unauthorized viewers.

Other one-way devices that attach to the cable network also suffer from the same vulnerability. The new CableCARD device for digital television is an example of such a device that suffers the same susceptibility to unauthorized redirection. Two-way devices, such as existing digital cable decoders for premium services, are less likely to suffer from this issue because there are means to electronically detect the location of these devices through headend interrogation and response, with the time delay to respond being measured to determine the cable distance to the device. In such an application, the response time values for two devices assigned to the same address can be compared for similarity and physical proximity inferred. While such a method provides some protection from unauthorized relocation, it will be shown later that it is an inferior method, suffering from poor resolution and other problems.

Sony’s has developed technology to address the issue of detecting unauthorized relocation of one-way equipment consigned to the subscriber by the cable operator. The method in which it is accomplished uses resources already available in the appliance and adds no additional hardware cost to the product.

The Network Interface for Digital Cable

Regardless of the end use of a particular device, all appliances attached to the digital cable network share a common front-end topology. The elements that make up the network interface are available from a number of different manufacturers and may be offered in different configurations featuring flexibility, integration with other elements, support of multiple interfaces, etc. to serve as the differentiation between products.

The typical network interface is shown in Figure 1. The cable network interface consists of two major sub elements, the RF tuner and the QAM demodulator. The function of the RF tuner is to receive all signals on the digital cable system and to exclude all but one desired RF channel, containing the digital service of interest. The method used to select the desired channel is called heterodyning and this process is used to convert an entire block of incoming signals to a lower intermediate frequency (IF), with the signal of interest centered on a fixed, constant value, which is passed through a fixed, narrow filter to eliminate the unwanted carriers. The QAM demodulator processes the tuner's IF output, converting it to an error free digital stream of MPEG transport data carrying the compressed audio and video services.

Inside the RF tuner the local oscillator, controlled by the host processor, varies in frequency such that the nonlinear combination of the local oscillator signal and the incoming spectrum from the cable network inside the mixer results in the signal of interest emerging from the mixer centered at the fixed, lower intermediate frequency. The IF typically might be selected to be a value such as 44MHz. The input filter eliminates extraneous signals outside the

range of valid cable audio/video services (54MHz to 863MHz) and the variable gain amplifier (VGA) is automatically adjusted so that the RF signals passing through the tuner and demodulator remain at optimum levels at all times. The final stage of the RF tuner is the surface acoustic wave (SAW) filter, which is an electromechanical device designed to only let a small band of signals centered at the IF value pass and all other RF energy to be heavily attenuated. The SAW only passes a standard 6 MHz wide channel and effectively rejects all others. The signal that emerges from the tuner is therefore only the channel carrying the service of interest and it has been downconverted to a fixed, standard (IF) frequency for processing by the QAM demodulator.

The QAM demodulator receives the incoming 6 MHz wide signal at the intermediate frequency, typically 44 MHz, and again amplifies it to a constant and optimum level through a second variable gain amplifier. The VGA is automatically adjusted by a closed control loop within the QAM demodulator. The signal is then processed by an analog to digital converter (ADC), which converts the incoming stream of time-varying voltages to a serial stream of binary bits representing the voltage levels of the signal at discrete time intervals. The ADC typically has 10 or more bits of resolution.

The digital stream is then split into two components, the in-phase component (I) and the out-of-phase component (Q). The Q term is used because the signal is in quadrature with respect to the I signal, meaning it is shifted 90° in phase. Phase separation occurs simultaneously with down conversion to a baseband signal, where the lowest frequency is 0Hz (DC) and highest frequency 6 MHz. This is in contrast to the incoming 44MHz IF signal, which has its content symmetrically

centered ± 3 MHz about the IF signal. The downconversion is accomplished through the use of a balanced mixer and the I-Q separation occurs because one of the two halves of the balanced mixer has a local oscillator signal output that is shifted 90° in phase relative to the signal applied to the other half of the balanced mixer. The outputs of the balanced mixer, I & Q, are then passed through identical channel filters that provide the appropriate shaping and attenuation of undesired processing artifacts occurring above the 6 MHz passband.

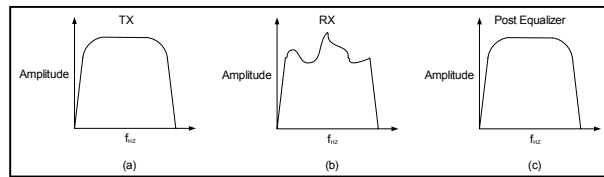


Figure 2. Digital cable channel spectrum

Next, an adaptive equalizer is applied to the outputs of the channel filters. The adaptive equalizer is an automatically self-varying digital filter network that continuously alters its filter characteristic (shape). Its purpose is to compensate automatically for echoes, reflections, dispersion, tilt, intersymbol interference and other distortions that alter the signal from its ideal, original form (Figure 2a) as it is carried by the cable operator's hybrid fiber-coax distribution network to the receiving device (Figure 2b), possibly over very long distances. By approaching the ideal of a matched filter, waveforms distorted through the communication path to the subscriber can be recovered and the data error rates for transmitted data reaching the phase recovery element (derotator) significantly reduced. This allows the system to operate successfully under non-ideal conditions, which are typical of real world applications.

The details of how the adaptive equalizer is realized differ between different QAM demodulator manufacturers. The general

architecture is common between them and takes the form of a classic feed forward/feedback digital filter. A typical digital filter for such a purpose is shown in Figure 3.

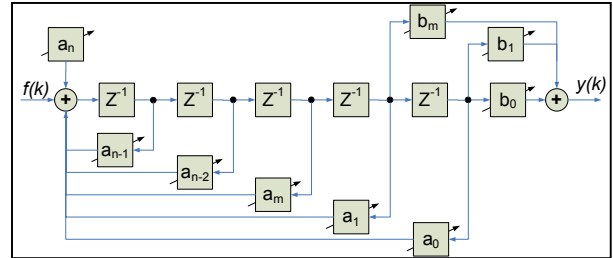


Figure 3. Adaptive equalizer

The structure of the filter is centered about a cascaded chain of delay stages (Z^{-1}), where the discrete time samples of the voltages seen at the demodulator input, converted to binary digital from by the ADC, are successively stored. The output of each delay stage tap, in addition to feeding the next cascade, may be fed back to the input or fed forward to the output. The tap feedback may be in conjunction with feed forward and either path may be employed exclusively on a tap-by-tap basis. Each feedback or feed forward path has associated with it an independent coefficient term. This term (a & b blocks in Figure 2) may provide amplification or attenuation of the tap output, depending upon the value of the coefficient. Because the equalizer is adaptive, the coefficients dynamically change under the control of a microprocessor or state machine. The values are varied based upon the characteristics of the equalizer output, as seen by the next processing stage, phase recovery. Typically a least mean square (LMS) algorithm is used to vary the tap values and converge upon the optimal solution. Adaptive equalizers in QAM demodulators vary in implementation between manufacturers. One design may have a total of 22 taps, where another may have a total of 40 taps – 16 feed forward and 24 feedback.

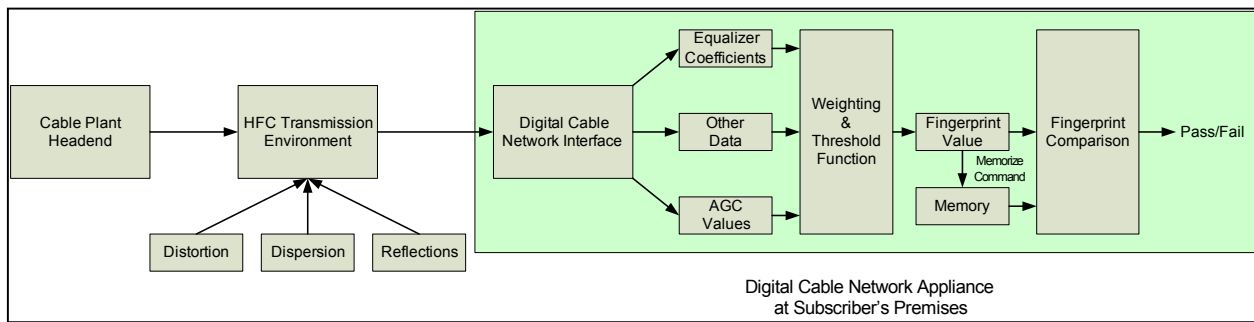


Figure 4. RF Fingerprint system

The output of the adaptive equalizer is then processed by the phase recovery block, also known as a detector or derotator. The purpose of the detector is to decode the combination of I and Q signals into a single data stream. The detector is able to expand the incoming data streams by a factor of $\log_2(\text{Modulation Order})$. This expansion is a factor of 6 for 64-QAM and 8 for 256-QAM, the two typical forms transmitted in digital cable. This expansion is the reason high transport data rates can be efficiently carried in relatively low spectrum bandwidths that appear to violate the Nyquist criterion. The coefficient values of the adaptive equalizer and the frequency setting of the QAM modulator local oscillator are both controlled by a microprocessor or state machine based upon the success of the detector to “lock” i.e. to recover valid data.

The last processing stage, the forward error corrector (FEC), applies a variety of algorithms to the raw recovered digital cable data stream to reduce the likelihood that any of the data has been corrupted in addition to formatting it appropriately for recovery of video and audio services as an MPEG transport stream. It is in this stage that de-interleaving, Viterbi (trellis) decoding, de-randomization, Reed-Solomon error correction and MPEG formatting occur. Some overhead data unique to the operation of these stages are removed from the stream so that the final MPEG transport emerging from the demodulator is identical in form, content and data rate to what the cable

operator inserted into the corresponding QAM modulator at the headend for transmission.

Further processing is done to decrypt, demultiplex, decompress and convert the content to a form suitable for display on a television. These steps, while vital to the proper function of a digital cable appliance, are beyond the scope of this document.

SONY'S “RF FINGERPRINT” TECHNOLOGY

The ability to detect changes in location of a one-way digital cable receiving device is based in large part upon the adaptive equalizer. The equalizer, as indicated, acts as a matched filter to the communications channel. As a result, the values contained within the equalizer's coefficients can be mathematically manipulated to show the transfer function of the communications channel that influences signals passing through it. Stated differently, the values of the coefficients, taken as a set, represent at a specific point in time the sum total knowledge of all mismatches, reflections, phase variations, gain variations, echoes and other perturbations of the transmission media upon the transmitted signal. The fact that the QAM demodulator is able to achieve and maintain signal lock under a given environment validates that the state of the equalizer at that time is such that it accurately reflects the knowledge of the plant's effect upon the system so it can

negate those effects and lock successfully. The tolerance to a suboptimal equalizer configuration is low, given the small vector error radii for either the QAM-64 or QAM-256 formats used in digital cable. The vector error radius is the composite of effects due to both amplitude and phase distortions upon a received signal.

Since the filter coefficient set is directly representative of the transmission environment, it responds dynamically to any changes in that environment. The low order feedback taps are most affected by high frequency trends, such as impedance variations at the connection or connector on the back of the appliance, reflections within the cable from the house splitter(s), etc. The middle taps are more predominantly affected by variations in the characteristics of the cabling to the tap and distribution amplifier, while the highest order taps are sensitive to channel tilt, dispersion, etc. This data, when combined with the AGC information which indicates total gain required for a constant signal level input, provides the basis for a very characteristic fingerprint of the environment where a specific cable appliance is installed.

Research by Sony indicates that the equalizer is so sensitive to such changes that one can distinguish between the short cables coming from different ports of an RF splitter to a bank of attached digital cable appliances fed by a single common source. In this case, the devices were all within one meter of each other and had identical cable lengths, yet the values observed for each device were unique and over time were relatively invariant.

If we let an equalizer coefficient be represented by $a \pm jb$, then H_1 , the matrix of all equalizer coefficients representing the state of the system at one point in time, i , can be represented by:

$$\mathbf{H}_{1,i} = \begin{pmatrix} a_0 & b_0 \\ \vdots & \vdots \\ a_n & b_n \end{pmatrix} \quad (1)$$

Likewise, if we let the gain value of one of the multiple nested AGC loops be represented by k , then H_2 , the matrix of all AGC coefficients representing the state of the system at one point in time can be represented by:

$$\mathbf{H}_{2,i} = \begin{pmatrix} k_0 \\ \vdots \\ k_n \end{pmatrix} \quad (2)$$

If one were to capture the equalizer tap coefficients and AGC data from a digital cable appliance, then applying an algorithm to allow the weighted summation of the coefficients a weighting function, based upon the expected statistical variance, to create a single scalar, a unary value representing the unique “fingerprint” of the environment of the device could be expressed. The threshold and weighting functions could be made unique to a particular operator and are kept secret to reduce the likelihood of tampering.

The algorithm for these operations then looks like:

$$Fingerprint_i = Y(\mathbf{H}_{1,i}, \mathbf{H}_{2,i}) \quad (3)$$

This fingerprint value is evaluated and stored in the digital cable appliance memory upon receipt of a command message, such as an EMM, from the cable operator. The stored value should be secured through encryption and digitally signed to detect tampering.

The Fingerprint Algorithm

Most of our current research on the RF Fingerprint concept is focused upon refining the algorithm used to calculate the RF Fingerprint, specifically the coefficient values of the weighting matrix, based upon experimental observation of actual cable appliances *in-situ*. The generalized form of the algorithm to calculate the fingerprint value based upon one manufacturer's implementation of QAM demodulator is:

$$Fingerprint = \sum_{i=0}^{39} (Tap_{i,a} \bullet W_{i,a}) + \sum_{i=0}^{39} (Tap_{i,b} \bullet W_{i,b}) + SF \quad (4)$$

Where W represents a value in the weighting coefficient matrix associated with a particular equalizer tap and SF represents the equalizer scale factor that is used to normalize all equalizer tap coefficients.

In order to determine the appropriate values for the weighting coefficients, tools were created by the team to remotely collect equalizer coefficient and other data from digital cable appliances installed throughout

the Sony cable test network, an elaborate system serving the entire San Diego corporate campus closely emulating a commercial cable television network. Devices were sampled at three minute intervals and remotely tuned to services on multiple frequencies. The data gathered was subsequently compiled into a large database for further analysis.

One of the first items confirmed was the repeatability of the equalizer configuration for a given subscriber drop and cable device. Repeatability was judged based upon the measurement of the standard deviation over a four day period. Figure 6 shows the results of one such data collection experiment where two different digital cable devices were attached to the test network for a statistically meaningful period. At the end of that period the two devices exchanged locations and data collection was restarted for the same period. As can be seen in the graph, not only is the data generally quite repeatable for a given device, but the dispersion tracks the location, not the device.

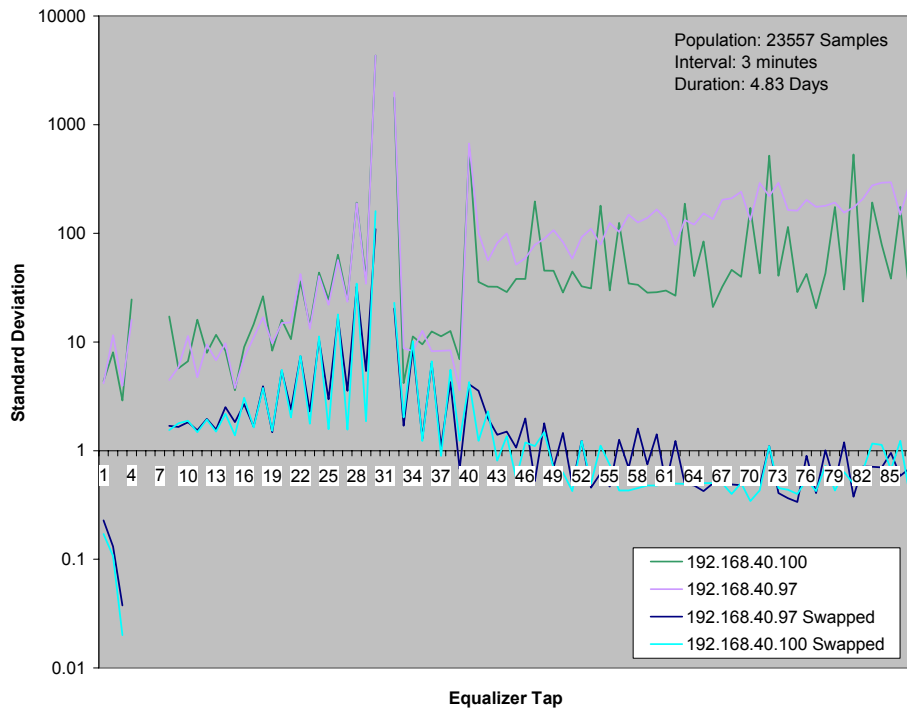


Figure 5. Positional uniqueness

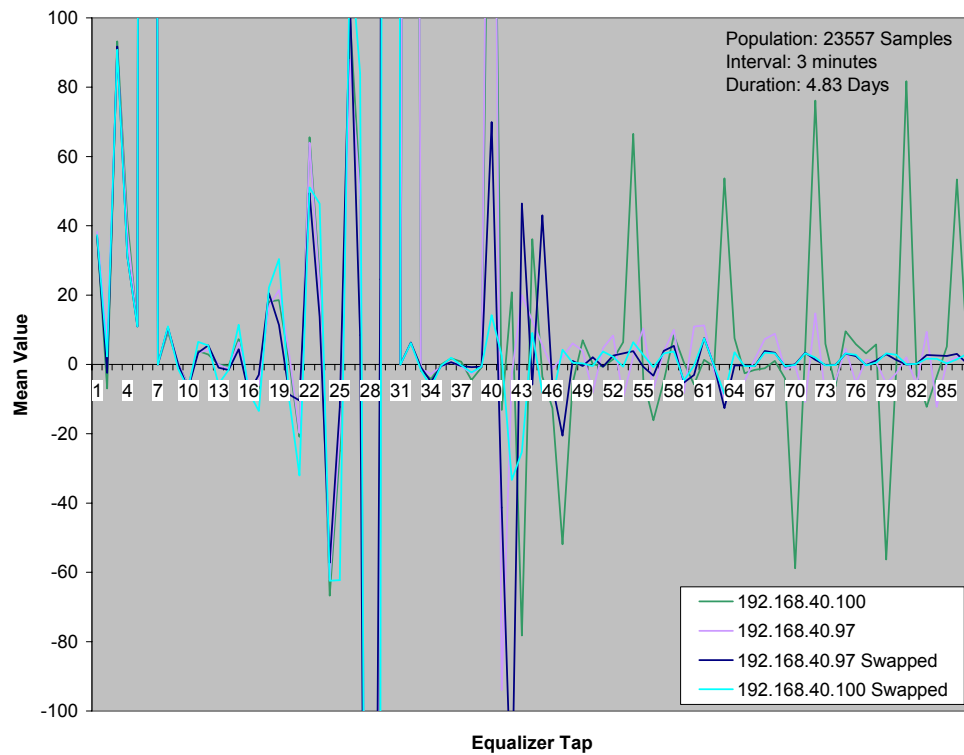


Figure 6. Equalizer coefficient repeatability

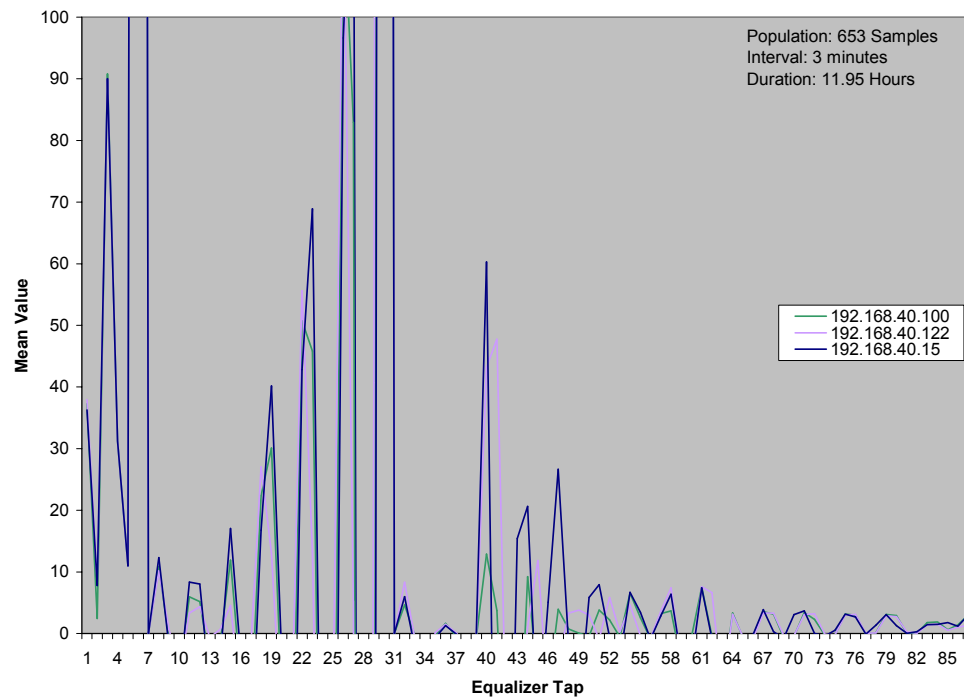


Figure 7. Uniqueness under worst case conditions

Using the data collected during the same experiment, Figure 7 shows that each of the four cases (two digital cable devices in two different locations) possesses a unique signature, allowing each to be distinguished from another device/location. The purpose of the weighting matrix is to selectively amplify those equalizer tap and AGC coefficients that express “uniqueness” terms and to attenuate those terms that contribute little in the context of distinguishing devices or tend to be unrepeatable.

In order to evaluate a worst case scenario, three cable appliances were connected through identical one meter long cables to a common splitter fed by the test network. The QAM performance of the three test units were monitored using four discrete frequencies at three minute intervals for a twelve hour period. The results of the investigation are shown in Figure 5. Confirmed was that even in an apparently identical RF environment and close physical proximity to ensure consistency of other environmental factors, the three different devices were distinguishable based solely upon equalizer and AGC coefficient values.

In practice, an installation having only one meter service drop lengths would seldom be seen and could actually be disregarded by the algorithm comparing fingerprint values to avoid false alarms due to product relocation from room to room within a home. However, it does prove that the concept is robust and applicable to high density dwellings such as apartments and dormitories where even two way methods, such as the one used in DOCSIS, fail. The DOCSIS method, for example, can only resolve location to within 63 meters, much more than the cable length differential between adjacent apartments. This is because DOCSIS and other proposed schemes use time delay measurement as the determining criteria. These other schemes

also require some form of two-way communication. The RF fingerprint scheme is immune to issues plaguing methods solely related to cable length and works in a purely unidirectional environment as well as bidirectional environments.

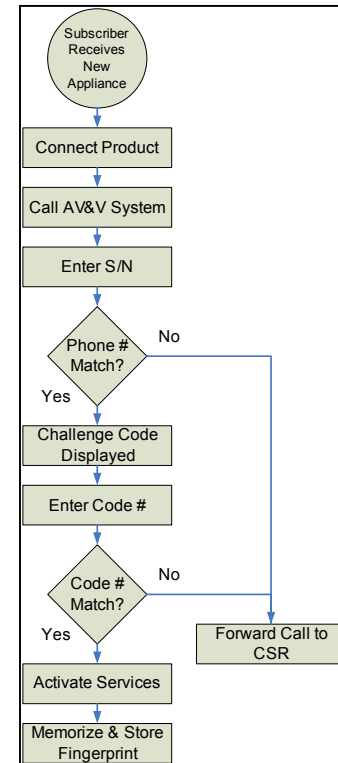


Figure 8. Digital cable appliance activation process

Operational Scenario

Regardless of the specific details associated with the implementation of the algorithm to calculate the fingerprint value, an operationally practical means of deploying a system employing RF fingerprint technology is necessary. Such a system must be automated to the greatest extent possible in order to reduce operating costs and to maximize flexibility. The practical implementation of a system employing RF fingerprinting is as follows:

1. A subscriber receives a fingerprint equipped appliance from the cable operator. The product contains labeling indicating “Call XXX-XXX-XXXX from your home phone after connecting the device to both the cable network and television for activation”. This is identical in concept to the process now followed for activation of home satellite television receivers and credit or ATM cards issued through mail by the major financial institutions.

2. The subscriber follows the instructions and calls the number on the appliance label after installation, as indicated. An automated validation and activation (AV&V) system at the cable operator receives the call and prompts the subscriber to enter the serial number of the cable appliance using the keypad on the telephone and to press the “#” key upon completion.

3. Upon receiving the “#” key input, the AV&V system confirms the validity of the entered appliance serial number. The system then looks through its subscriber database and finds the record for the subscriber issued the appliance having the entered serial number. It then reads the subscriber’s home phone number from the database record. Using Automated Number Identification (ANI), a non-maskable form of caller identification used for logging calls to toll-free telephone numbers (and 911 calls), the AV&V system then confirms a match between the number of record and if unsuccessful, refers the call to a customer service agent. This step validates that an authorized subscriber is attempting to activate the appliance issued to them by the cable operator.

4. If the ANI and phone number on record match, the AV&V system then sends a control message (EMM) to the appliance having the serial number the subscriber entered by phone. This EMM commands the appliance to display, on the subscriber’s television screen, a challenge number sequence contained within the EMM message and generated at random by the AV&V system. The AV&V system then instructs the subscriber to enter into the telephone, the number displayed on the screen using the keypad on the telephone and to press “#” key upon completion.

5. Upon receiving the “#” key input, the AV&V system confirms the validity of the entered challenge number and if unsuccessful, refers the call to a customer service agent. This step validates that the authorized subscriber is attempting to activate the appliance issued to them by the cable operator at the home on record.

6. If the challenge sequence is successful, the AV&V system sends another EMM to the now validated appliance, commanding it to perform two steps

A. Activate the services authorized for that subscriber

B. Calculate the RF fingerprint for the appliance at the present location and store it in persistent memory.

At periodic audit intervals, determined either by EMM or through self-initiation, which uses a timer resident in software, the appliance collects data and calculates an electronic fingerprint value, comparing it to the reference value stored in memory. If the calculated value is within predetermined

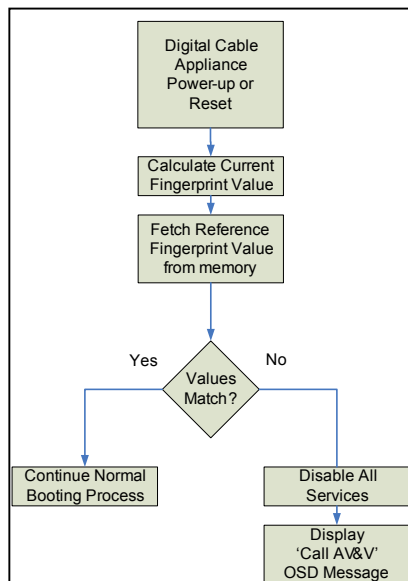


Figure 9. Digital cable appliance location audit process at reboot

limits, no further action is taken until the next audit period. If the new value for the fingerprint is sufficiently different from the stored reference value, then the reference value in memory is updated with the new reference value.

Whenever the appliance is rebooted or otherwise reset, signifying a lapse in network connectivity where the appliance may have been relocated without the authorization of the cable operator, the appliance collects data and calculates an electronic fingerprint value, comparing it to the reference value stored in memory. If the calculated value is within predetermined limits, the device continues the booting process and services are restored. If the match is unsuccessful, all television services are automatically self-deauthorized on the appliance, with an on-screen message generated and displayed on the subscriber's television screen indicating that the cable operator must be contacted at the AV&V telephone number contained within the

message for appliance reactivation. This message occurs because the appliance has determined that an unauthorized relocation has possibly occurred. When the subscriber calls the displayed telephone number, the AV&V process is executed and upon validation, the location of the device is re-evaluated.

CONCLUSION

RF fingerprinting technology is one element in the toolkit available to address issues of service and equipment theft in cable television systems. Implemented in a digital decoder and coupled with simple content encryption techniques, a complete solution providing both the quality of service seen in an all-digital network as well as the system security and compensation for services delivered previously encountered only in premium digital services. All this is possible in the lowest possible cost, one-way customer premises equipment. The implementation of RF fingerprinting does not add hardware cost to the cable device and can be implemented in any digital device attached to a cable television network.

On-going development of this technology continues to focus upon optimization of weighting matrices as the technology is matured. One major U.S. cable operator has already specified the inclusion of this technology in their current system upgrades to all-digital delivery and draft specifications of the management and control aspects have been recently completed. It is quite likely that within the next 18 months, a commercial example of RF fingerprinting technology will be available and in the hands of the public.

ACKNOWLEDGMENT

The author would like to thank Eric Holcomb and Aran Sadjja, both of Sony EMCSA in San Diego, for their significant contributions to the development of the RF fingerprint technology as well as the tools and systems implemented to perform automated data collection in order to validate the concept.

REFERENCES

1. ITU-T Rec. J.83:1997, *Digital multi-programme systems for television, sound and data services for cable distribution*. Geneva: International Telecommunication Union, Apr. 1997.
2. ISO/IEC 13818-1:2000, *Information Technology – Coding of moving pictures and associated audio – Part 1: Systems*. Geneva: International Organization for Standardization/International Electrotechnical Commission, Dec. 2000.
3. B. P. Lathi, *Signals, System and Controls*. New York: Harper & Row, 1974, pp. 207-214 & pp. 428-456.
4. A. Bruce Carlson, *Communication Systems*, 3rd ed. New York: Mc Graw-Hill, 1986, pp. 514-517 & pp. 550-554.
5. Edward A. Lee and David G. Messerschmitt, *Digital Communication*, 2nd ed. Boston: Kluwer Academic, 1994, pp. 442-550.
6. Richard E. Blahut, *Digital Transmission of Information*. New York: Addison-Wesley, 1990, pp. 159-170.

SEAMLESS MOBILITY BETWEEN HOME NETWORKS AND CABLE SERVICES

Jay Strater and Gordon Beacham
Motorola

Abstract

Cable operators are now offering landline voice service in addition to video and high-speed data. However, cellular service is quickly becoming another important service for cable operators to offer to ensure that they stay in a competitive position with telephone companies.

This paper provides an overview of seamless mobile communications and its potential use by cable operators. It includes a description of seamless mobile applications and their value, network approaches for seamless mobility, and considerations for power management, Quality of Service (QoS), security, and Network Address Translation (NAT) traversal. The paper concludes with a discussion of other seamless mobility concepts revolving around the connected home.

INTRODUCTION

Through a careful development and standardization processes culminating in Data Over Cable Service Interface Specification (DOCSIS) in the late 1990's, cable operators have very successfully deployed high-speed data as a key element of the broadband service package. Cable operators are leaders in residential broadband, and are now leveraging this penetration to provide landline voice service using Voice-over-IP technology (VoIP) to POTS phones.

Landline voice service provides the "third leg" of cable's current "triple play", completing a service mix consisting of data,

video, and voice. Nonetheless telephone companies are a competitive threat since teaming up with satellite-TV providers to offer their own voice, high-speed internet, and video service bundle. The three largest local telephone companies also own significant portions of the major wireless phone carriers, giving them a leg up on the cable providers by also offering cell phone service. Consequently, as enticing as landline voice service is, cellular service is quickly becoming another important service, a "quadruple play", for cable operators to offer to ensure that they stay in a competitive position.

Cable operators are interested in offering a mobile service with access to cellular and wireless local area network (WLAN) service using seamless mobility. In general terms seamless mobility is an approach that allows users to roam between application domains and communication networks without being aware of the underlying mechanisms that enable them to do so. This includes the scenario where a user moves between environments where different networking capabilities are present, but the network provides negotiation to allow for seamlessly transparent access. This differs from today's environment where handover between heterogeneous networks is not supported in most cases, and users are required to stop one communication service and initiate another between different networks.

Seamless mobile communications between cellular and WLAN environments refers to a service in which a user receives voice and data service on a dual-mode mobile handset

when inside or outside a WLAN.¹ Service within a WLAN is provided via VoIP over unlicensed WLAN spectrum and DOCSIS in the case of a cable plant. Service outside the WLAN is provided via circuit switch cellular technology (GSM or CDMA) or VoIP over cellular or a 3G data network. Moreover, a voice call or data connection is seamlessly handed off between WLAN and cellular network as a user moves between them.

This paper provides an overview of seamless mobile communications and its potential use by cable operators. It includes a description of seamless mobile applications and their value, network approaches for seamless mobility, and considerations for power management, QoS, security, and NAT traversal. The paper concludes with a discussion of other seamless mobility concepts revolving around the connected home.

SEAMLESS MOBILE COMMUNICATIONS APPLICATIONS AND VALUE PROPOSITION

As described previously, we refer to seamless mobile communications as a service in which a user receives voice and data service on a dual-mode mobile handset with cellular and WLAN capability. Moreover, service in the WLAN is provided via VoIP over Wi-Fi and carried via DOCSIS in a cable plant. Service outside the WLAN is provided via cellular technology such as GSM or CDMA. A voice call or data connection is seamlessly handed off between WLAN and cellular networks as a user moves between them.

Seamless mobile communications may apply to residential or enterprise applications. Cable operators are focused on residential seamless communications initially but see

¹ A dual mode handset here refers to a mobile handset with functionality for cellular and WLAN operation.

value in an enterprise application longer term. A few differences exist between residential and enterprise applications. These include the number of access points, the call control technology and features, and the level of required security.

Residential applications typically require only one WLAN access point (AP) in a user's home whereas enterprise applications require multiple WLAN AP in a user's work location. Having multiple APs adds to the complexity of a seamless mobile communications solution by requiring micro-mobility management for call handover between APs in addition to macro-mobility handoff between WLAN and cellular service. In addition residential applications have phone service controlled from a central office (circuit or packet switched) and CLASS5 features offered, whereas enterprises typically have phone service controlled by a PBX (often an IP PBX using SIP or H.323 signaling protocols) along with enterprise features.

What are the benefits of seamless mobile communications? For residential users, the benefits may include:

- Reduced cellular bill resulting from off-loading the cellular air-interface when calls are made from the mobile handset in the home's WLAN²
- Improved in-home coverage and reliability which is often limited with cellular service
- Wireline audio quality because of a higher-rate codec thanks to WLAN and broadband connections

² Cell phone contracts typically bill for a bulk number of minutes and residential VoIP service is typically an all-you-can-eat service. By switching your cell phone call to a Wi-Fi call, the minutes should now fall in the all-you-can-eat category.

- Convenience of a single mobile number and voice mail service, whether inside or outside the home
- Mobile and landline voice service interworking, e.g. allowing for a shared “family” number as well as “individual” mobile and landline numbers

Studies indicate that users will still want access to their existing landline telephone service in addition to a mobile handset. Consequently, residential applications should also allow for a mobile handset and landline phone to inter-work. In the example noted above a “family” number implies a user selected mix of mobile and landline phones, allowing for group ringing and possibly line-extension service. By contrast an “individual” number implies a unique number for mobile handsets and landlines. Having “family” and “individual” numbers requires that distinctive ringing be supported.

For enterprise users, the benefits of seamless mobile communications may include:

- Cost savings from eliminating use of cellular calls in the enterprise, assuming cellular service is available and otherwise utilized in the enterprise
- Operational cost benefits from carrying VoIP and data on a single network; although this may already be realized if an IP PBX service is in place
- Wireline audio quality because of a higher-rate codec thanks to WLAN and broadband connection
- Convenience and productivity benefit of using a single device for all communications, whether inside or outside the office, whether at or away from ones desk

In addition to the benefits outlined above, users may receive advanced services and, through use of a common handset, consistent services inside and outside the enterprise and home.

For wireless carriers, the benefit of seamless mobile communications may include:

- Cellular capacity relief by off-loading of air-interface when user is in a WLAN at home or in the office
- Ability to offer residential phone service, in the case of independent wireless carriers (those without a wireline carrier affiliation)
- Increased customer satisfaction from improved in-home coverage, a key user network quality metric
- Improved customer retention (reduced churn) through unique value added services

For cable operators, the benefit of seamless mobile communications may include:

- A “quadruple play” service offering of data, video, voice, and wireless as in the case of cable operators, allowing added service bundling
- Greater pricing flexibility resulting from increased service bundling and migration of customers to higher revenue / margin wireless offerings
- Increased leverage of existing broadband infrastructure
- Ability to offer superior QoS over cable infrastructure (via DOCSIS 1.1 or higher QoS)
- Increased customer retention through bundling and possible value added services, including interworking of mobile and landline phone services

The presumption is that by bundling services, a broadband carrier will present a customer with only one bill, and offer some service discounting compared with offering services individually.

Finally for vendors, the benefit of seamless mobile communications may include:

- Increased handset, modem, and access point /gateway sales
- Increased core network equipment sales
- New opportunities for integration and deployment services

In summary, seamless mobile communications appears to offer a win-win solution for all stake-holders. And as stated in the introduction, having a “quadruple play” is quickly becoming a service cable operators will need to stay competitive.

NETWORK APPROACHES FOR SEAMLESS MOBILE COMMUNICATIONS

An important aspect of any seamless mobility solution is providing an easy-to-use service that is transparent to the user. Ease of use and seamless service functionality can be achieved by including intelligence in the network or in end user devices. The most likely point for intelligence to be implemented is in the network with varying levels of support in the end user devices.

We know of three fundamentally different network approaches to providing mobile seamless communications. They include call forwarding, Unlicensed Mobile Access (UMA), and Third Generation Partnership Project (3GPP) IP Multimedia Subsystem (IMS) technologies.

Call Forwarding

In a call forwarding approach, a specialized mobility management element is integrated (logically or functionally) with a call management server (CMS) serving landline phones and mobile/dual-mode handsets in a WLAN. For residential applications the CMS would reside in a central office whereas for enterprise applications the CMS may reside in the enterprise as an IP PBX. In either case there is no need for changes at a cellular Mobile Switching Center (MSC).

A phone number is assigned at the CMS to each dual-mode handset as well as landline phones served by the CMS. Calls destined to a dual-mode handset’s CMS number when the handset is inside a WLAN are served directly by the CMS. Calls destined to the dual-mode handset’s CMS number when the handset is outside the WLAN (in the handset’s cellular WAN) are forwarded by the CMS to the handset’s cellular number. Seamless handover between (into or out of) WLAN and cellular WAN is provided for any calls to the dual-mode handset’s CMS number, and any calls initiated by the dual-mode handset from inside the WLAN or to a phone controlled by the handset’s CMS.

For single number reachability and seamless mobility to be supported as outlined above, all calls must pass through the CMS/mobility management element(s). Consequently the handset’s CMS number must be the only number advertised to friends and colleagues.

The following figure illustrates an example of a call-forwarding network architecture. VoIP signaling and bearer transport protocols are shown, data is not.

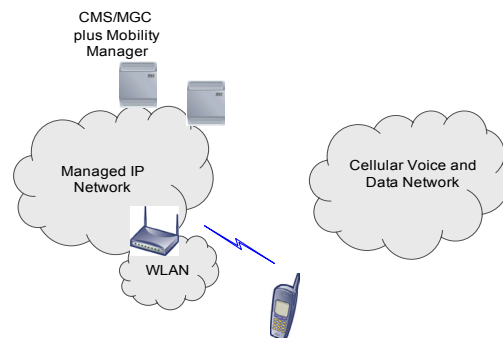


Figure 1. Example Network Architecture for Call-Forwarding Approach

The call-forwarding approach has the benefit of making maximum use of existing CMS and media gateway resources if present. As such new integration requirements with operator back-office systems may be

minimized and use of common features for landline phones and mobile handsets may be maximized. The call forwarding approach has the disadvantage of imposing inefficient routes via dog-legged calls that must always be anchored at a CMS, even if the CMS is far from source and destination. This inefficiency may prove costly from a PSTN interconnect perspective and also from a transcoding, performance perspective, when transcoding is required.³ Also support for seamless mobile handover from cellular WAN to WLAN is not possible when calls are initiated by a dual-mode handset in the cellular WAN to a phone outside the control of the dual-mode handset's CMS. Finally seamless mobility call forwarding approaches are not currently standardized, raising the concern over single supplier cost premiums and limited product availability.⁴

Unlicensed Mobile Access

With UMA technology, GSM and GPRS signaling and RTP and IP data traffic are carried over an IPsec tunnel between a dual-mode handset and UMA network controller (UNC) when the handset is in a WLAN. The UNC emulates a GSM base station controller (BSC), providing a gateway between WLAN and cellular WAN signaling and traffic. GSM signaling is used to signal when a user moves between a UNC and a BSC for mobility management and handover purposes. The cellular number maintained at the cellular carrier's home location registrar (HLR) identifies the user.

³ Transcoding between WLAN and PSTN can be avoided by utilizing G.711 while in the WLAN.

⁴ Even if signaling from handset to mobility manager is SIP based, unique SIP extensions are likely required to support specific mobility and handoff signaling requirements. To date no call forwarding approaches we know of have been standardized. However Motorola, Avaya and Proxim have developed an enterprise solution that they are working to move into an industry standard.

The following figure illustrates an example UMA network architecture. VoIP signaling and bearer transport protocols are shown, data is not.

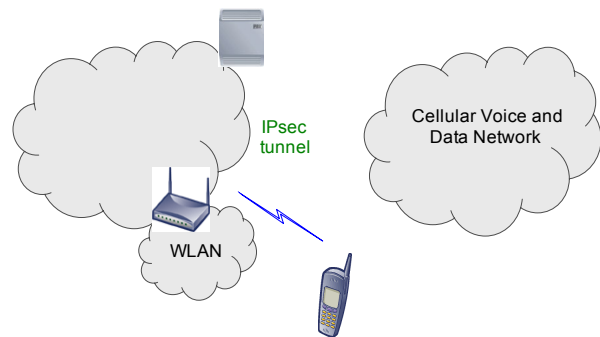


Figure 2. Example UMA Network Architecture

The UMA approach has the benefit of requiring little change to an existing GSM cellular infrastructure, only requiring the addition of a UNC. As such the approach offers a quick time-to-market solution for cellular operators. It also has specifications that are governed by a consortium of cellular operators and vendors. The UMA approach has the disadvantage of being limited to GSM systems. And because it utilizes existing GSM infrastructure it is limited in terms of offering new, advanced IP-based services (limited to what an existing MSC offer). Moreover a major investment in current-technology GSM infrastructure including MSCs would be required for those operators without GSM infrastructure (most MSOs) to deploy a UMA solution.⁵

IP Multimedia Subsystem

⁵ Otherwise, an unorthodox business arrangement would be required in which a cellular operator Mobile Switching Center (MSC) allows access from a 3rd party operator's UNC. In any case, cable operators could still receive revenue from a UMA solution offered by cellular operators over a cable operator's infrastructure. In this case the cable operator could provide the cellular operator QoS on the DOCSIS network and charge for it accordingly. Of course cellular operators can also use the DOCSIS network without permission or knowledge of the cable operator and offer a best effort VoIP service.

With IMS technology, VoIP is provided via SIP signaling and RTP traffic between a dual-mode handset and call control elements in an IMS core network when the dual-mode handset is in a WLAN. The same applies if a dual-mode handset is in a 3G cellular network that supports VoIP over its data network. When a handset is in a circuit-switched cellular network, circuit-switched signaling and voice are converted to SIP signaling and RTP at an IMS signaling and media gateway respectively. To neighboring MSCs in a circuit-switched cellular WAN, the IMS core looks like another MSC. Yet all signaling within the IMS core network remains SIP based with WLAN and cellular call mobility controlled from the core.⁶ In essence, the IMS core provides WLAN and cellular network-agnostic access control with centralized application services.

A cellular number is maintained at the home location registrar (HLR) associated with the IMS core network and identifies the user. However additional numbers at the IMS core could be used for identifying a dual-mode handset when in the WLAN. This allows flexibility to support “family” and “individual” phone numbers when the handset is in the WLAN.

Note that IMS, which is developed by the 3GPP forum, applies to GSM-based cellular network technologies. In contrast, CDMA-based cellular network technologies have an IP core solution like IMS as directed by 3GPP2 Multimedia Domain (MMD) forum. Therefore, we refer to IMS as a cellular core technology for supporting any cellular network technology. Note also that seamless handoff between cellular networks and/or WLAN has not been specified as yet. Its specification is anticipated as another SIP application serving an IMS serving call session control function (S-CSCF).

⁶ 3GPP SIP header extensions are used in IMS SIP signaling.

The following figure illustrates an example IMS network architecture. VoIP signaling and bearer transport protocols are shown, data is not.⁷

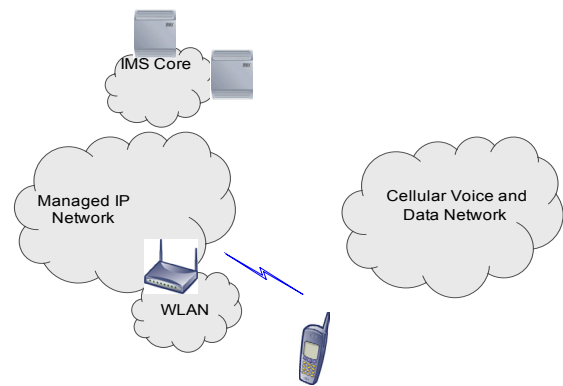


Figure 3. Example IMS Network Architecture

Compared to UMA, the IMS approach has the benefit of being RAN agnostic and facilitating rapid deployment of enhanced packet-switched services. It also has the benefit of allowing a cable operator the flexibility of controlling a dual-mode handset when on a WLAN connected to their broadband network. Compared to call forwarding, the IMS approach has the benefit of being governed by a set of standards (developed by 3GPP) and avoiding route, functionality and potential performance issues of the call forwarding approach. A limitation of IMS is that it does not make use of existing Call

⁷ Of the IMS signaling elements: the HSS provides Home Subscriber Server with AAA and Databases. The Application Servers (AS) include a Session Initiation Protocol (SIP) AS, an Open Service Access (OSA) Service Capability Server (SCS) & OSA AS, and an AIN Interworking Server. The CSCF is the Call Session Control Function which has 3 flavors. Serving-CSCF provides session control for endpoint devices, Interrogating-CSCF is an entry point to IMS from other networks, and Proxy-CSCF is an entry point to IMS for devices. The BGCF is the Breakout Gateway Control Function which selects networks to use for PSTN/PLMN interworking. The MGCF is the Media Gateway Control Function which controls the MGW. The MRFC is the Multimedia Resource Function Controller which controls MRFP (media server). And the PDF is the Policy Decision Function which authorizes QoS requests.

Management Servers (CMS) and media gateway resources for VoIP-based landline phone service. Dual-mode handset and landline inter-working is expected to be the subject for Phase 2 of the PacketCable 2.0 cellular integration initiative.

In summary, of the three alternatives discussed, IMS appears to offer the most promising solution to cable operators from a business, functional, cost, and performance perspective.

POWER MANAGEMENT, SECURITY, QOS, AND NAT TRAVERSAL CONSIDERATIONS

Aside from providing seamless roaming between WLAN and cellular WAN and enhanced services, seamless mobile communications solutions must address several critical factors in order to provide a viable service. These include handset power management, security, QoS, and NAT transversal for operation from a WLAN.

Power Management

Power management is required for a dual-mode handset in a WLAN to have active and idle talk time that is comparable to that in typical cellular service. Without power management active talk time may be reduced from hours to minutes and idle talk time may be reduced from tens of hours to single-digit hours.

Through the development of an enterprise seamless mobility solution, Motorola identified several factors as necessary for improving battery life in a dual-mode handset. These include choosing a low-power 802.11 chipset in the handset, running the handset with one radio at a time (cellular WAN or WLAN), providing fast switching technology to select WLAN/WAN, and implementing a new 802.11 power management approach for VoIP traffic. In the latter case, legacy 802.11 power management alone was found to be

inadequate. To understand what new 802.11 power management approach was selected for VoIP traffic, it is first necessary to understand why legacy 802.11 power management was found inadequate.

In 802.11, legacy power management allows a station (dual-mode handset in our application) to receive data from an AP at infrequent intervals and only when data is available at the AP for delivery. AP data transfer may be scheduled to occur at Delivery Traffic Indication Message (DTIM) beacons, spaced say 300ms apart. Utilizing this timing, a station need only wake its processor infrequently to listen to selected AP beacons and poll the AP for its data if the AP indicates its data is available (via Traffic Indication Map).⁸ However, VoIP traffic sent between station and AP during a VoIP call requires a short packet exchange intervals, say 20ms. Consequently, if legacy 802.11 power management were used it would need to have shortened beacon intervals of say 20ms to accommodate the VoIP traffic. In that case, processor sleep time would be very limited, particularly during idle call periods.

The new 802.11 power management approach is a method for allowing periodic frame exchange between a station and AP during a VoIP call while legacy power management is still utilized for data frame exchange. A station activates the new power management approach when it negotiates an admitted VoIP traffic stream (or any periodic frame transfer stream) with its AP. An AP then buffers frames from the VoIP traffic stream destined to the station and, when a station wakes up at a designated VoIP frame interval (presumably set to the VoIP packetization period), it sends a VoIP frame with an implicit "poll" request to its AP. The

⁸ An AP responds to a poll with a buffered frame. The handset must poll for more buffered data if more exists. The AP will also send broadcast and multicast data right after DTIM beacons.

AP responds to a poll request with buffered VoIP stream frame(s).

This new power management approach has been pushed into the 802.11e standard where it is referred to as Unscheduled Automatic Power Save Delivery (U-APSD). U-APSD may be enabled when the U-APSD station issues a bi-directional Traffic Specification (T-SPEC) to its AP at the start of a VoIP call⁹. The AP buffers frames matching the Access Category (AC) for the T-SPEC according to the method described previously. The following figure illustrates U-APSD from the perspective of AP and station (subscriber unit).

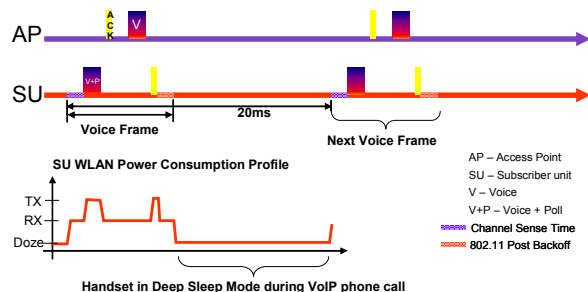


Figure 4. U-APSD Operation

Security

Security for a dual-mode handset operating from a WLAN is necessary to prevent user traffic (VoIP and data) from being denied, stolen, or eaves-dropped. The level of security offered should be commensurate with the level of threat and the impact of a successful security attack.

Security should be provided on the WLAN between handset and AP and end-to-end between handset and network infrastructure. For Wi-Fi access security, several standard approaches are possible. For residential applications, wireless protected access with pre-shared key (WPA-PSK) is a logical choice. WPA-PSK is based on the 802.11i

draft 3 for authentication and encryption. Using WPA-PSK, mutual authentication, key verification and derivation between the station/handset and AP is established using an 802.1x handshake. This is illustrated in the following figure.

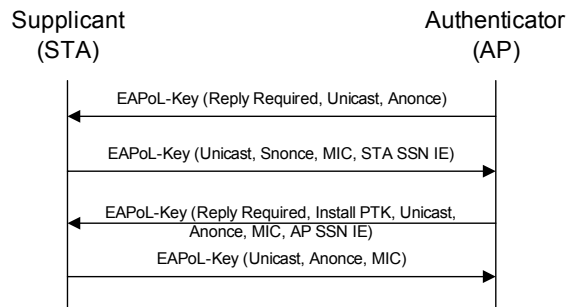


Figure 5. 802.11i WPA-PSK Example with Supplicant and Authenticator

After successful authentication, the Temporal Key Integrity Protocol (TKIP) encryption algorithm is used to encrypt session traffic. WPA2-PSK is another logical choice for residential applications. WPA2-PSK is an evolution of WPA-PSK, based on the final ratified 802.11i standard and includes use of the Advanced Encryption Standard (AES) encryption algorithm instead of TKIP.

For enterprise applications, WPA or WPA2 using an operator selected Extensible Authentication Protocol (EAP) method is appropriate. EAP methods may include mutual authentication schemes based on X.509 certificates (i.e., EAP-Transport Layer Security (TLS)), or other alternatives such as EAP-Tunnel TLS (TTLS), EAP-Subscriber Identity Module (SIM), etc. EAP message exchanges are illustrated in the following figure.

⁹ 802.11e draft 12 also provides a mechanism for a station to enable U-APSD in its reassociation request frame, without submitting a T-SPEC.

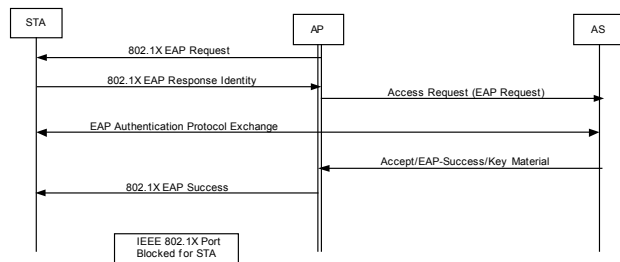


Figure 6. 802.11i Example with Supplicant, Authenticator, and Authentication Server

For end-to-end security, handset authentication and authorization is necessary to mitigate theft of service. In addition an IPsec tunnel may be utilized to protect VoIP, data, and signaling traffic. UMA and 3GPP IMS standards specify use of an IPsec tunnel and EAP authentication. In the case of UMA, an IPsec tunnel is established between a handset and UNC. The tunnel follows authentication between handset and UNC secure gateway and AAA server. Authentication is performed using EAP-SIM within IKEv2. IKEv2 is used to authenticate a UNC to a handset and EAP SIM is used to authenticate a handset to a network Authentication Authorization and Accounting (AAA) server, as relayed by the UNC's security gateway. UMA requires authentication and establishment of an IPsec tunnel with at least two UNC's as part of handset discovery and registration.¹⁰

In the case of IMS, an IPsec tunnel is established between a handset and Packet

Data Gateway (PDG). As part of the tunnel establishment, the PDG contacts a 3GPP AAA Server in the Home Public Land Mobile Network (HPLMN) for authorization of the mobile.¹¹ IMS may require several IPsec tunnel attempts, starting with PDG associated with a handset's Visiting PLMN and progressing to PDG associated with its HPLMN.¹² Authentication between handset and AAA server uses either EAP-SIM or EAP-Authentication and Key Agreement (AKA).

Note that enterprise solutions may also utilize separate VPN security for remote data access.

QoS

QoS for VoIP calls from a dual-mode handset when operating in a WLAN should be comparable to landline-like voice service to differentiate service from competitive VoIP service providers. QoS should be provided over the WLAN between handset and AP, and end-to-end between handset and network infrastructure.

The 802.11e standard provides QoS methods for a Wi-Fi network. It includes specifications for prioritized and parameterized QoS in two basic access approaches: Enhanced Distributed Channel Access (EDCA) service and Hybrid Coordination Function Controlled Channel Access (HCCA). EDCA is a prioritized CSMA/CA access mechanism offering four differentiated service Access Categories (AC). EDCA also offers a level of parameterized

¹⁰ UMA clients are provisioned with or derive the FQDN for their provisioning UNC and UNC SGW. After resolving these addresses, a UMA client sets up a secure tunnel with the provisioning UNC to discover the FQDN or IP address of its default UNC and UNC Security Gateway (SGW). The UMA client then sets up a secure tunnel to its default UNC and attempts to register. The UMA client's default UNC may deflect the UMA client to an alternate serving UNC. In this case the UNC client will establish a secure tunnel this UNC and register.

¹¹ The service authorization decision is made by the 3GPP AAA Server based on subscription information retrieved from the IMS HSS/HLR.

¹² A handset will initially construct an FQDN based on its Visitor Public Land Mobile Network (VPLMN) and use it to discover (via DNS) a set of visited-network PDG IP addresses to attempt IPsec tunnel establishment with. Should these addresses not be reachable or the AAA server reject access to them, the handset will then construct a new FQDN based on its HPLMN and attempt IPsec tunnel establishment with a home-network PDG.

QoS for time sensitive traffic flows such as VoIP. A T-SPEC with AP admission control is used to establish a traffic stream. HCCA is a more complex approach involving a mix of contention and contention free periods (CFP). It provides T-SPEC parameterized QoS for time sensitive traffic using scheduled transmission opportunities in the CFP.

Several subsets of the 802.11e standard were identified by the Wi-Fi Alliance in order to expedite interoperability testing for WLAN QoS. Wireless Multimedia Extension (WME) features and more recently the Wireless Multi Media (WMM) features were chosen based on the EDCA. Compliance to these tests is mandatory for many solutions, making EDCA commonly utilized. With EDCA, VoIP calls are typically assigned a T-SPEC with QoS parameters that limit delay and jitter and with AC value set to VO, the highest level.

For end-to-end QoS, IP differentiated services per RFC2474 (DiffServ) should be considered. DiffServ sets packet priority via a Differentiated Services Code Point (DSCP) field of an IP header. For upstream traffic from dual-mode handsets in a Wi-Fi network a DSCP assignment could be used to classify a VoIP packet onto an 802.11e VoIP T-SPEC with AC_VO, established at the start of the call. It could also be used to classify a VoIP packet onto a DOCSIS Unsolicited Grant Service (UGS) service flow in the case of cable modem service, also established at the start of the call. In this case, PacketCable Multi-Media (PCMM) is needed for establishing a UGS service flow for the upstream VoIP call.¹³

¹³ PCMM QoS is initiated by an Application Manager which, for a seamless mobility solution, would need to be affiliated with an IP PBX or CMS in a call forwarding solution, the UNC in a UMA solution, or an S-CSCF in an IMS solution.

In the case of downstream end-to-end traffic a reverse process would apply. A DSCP assignment at the other end of the VoIP call (terminating handset, MTA, or media gateway) could be used to classify a VoIP packet such that it receives proper DOCSIS downstream service flow assignment at the CMTS and 802.11e T-SPEC at a Wi-Fi AP. Downstream service flows would also have been established at the start of a VoIP call.¹⁴

The 802.1d or 802.1p user priority should also be assigned for priority queuing at an Ethernet switch. The following table provides a listing of typical assignments for 802.1d, 802.11e AC, and RFC2474 DSCP.

Table 1. Typical Relationships for QoS Assignments

8021D User Priority	80211e Access Category	RFC2474 DSCP Bits	Recommended Traffic Category
1, 2	AC_BK	001xxx or 010xxx	Background
0, 3	AC_BE	000xxx or 011xxx	Best Effort
4, 5	AC_VI	100xxx or 101xxx	Video
6, 7	AC_VO	110xxx or 111xxx	Voice

Note that there are various techniques that can be used for associating signaling and media to an 802.11e AC. In one case a station/handset could have a default AC method for signaling and media. In another case, a station could use UPnP to discover a UPnP capable AP's QoS capabilities and adjust its signaling and media ACs accordingly.

NAT Traversal

Dual-mode handset traffic (voice or data) should be allowed to pass through any residential gateway for a seamless solution to be widely deployed. This includes gateways with network address and port translation (NAPT).

¹⁴ In the case of a DOCSIS downstream, the service flow would need a minimum reserved rate to be assigned via PCMM.

Two issues concern setting up and passing VoIP traffic through a NATing gateway. The first issue is that protocols that embed local IP or port information in signaling messages (i.e., SDP in SIP) will not be reachable through a NATing gateway without a solution that provides the device behind the gateway (dual-mode handset) with information about its NATed IP address and port or without a network server that provides address translation from local to routable addresses.¹⁵ Simple Traversal of UDP through NAT (STUN) or Traversal Using Relay NAT (TURN) are two approaches that provide a client device behind the gateway with information about its NATed IP address and port.

There are four types of NAT implementations: full cone, restricted cone, port restricted cone, and symmetric as defined in [9]. STUN is a solution for full cone, restricted cone, or port restricted cone NAT implementations. However, in the case of a symmetric NAT implementation, STUN is not an effective solution because it will not work with connection oriented media using RTP. This is because a symmetric NAT will have different NAT mappings for each destination IP address and port. The symmetric NAT requires a VoIP client to send and receive RTP packets to and from the same IP address and port. For symmetric NAT implementations, TURN is a possible solution. TURN involves using an RTP relay server between communicating VoIP clients.

Unfortunately, the TURN solution has some disadvantages. It is processing intensive on the RTP relay server, raising issues of scalability, and it introduces

¹⁵ A gateway ALG may also be used but it is limited to scenarios in which signaling is not encrypted, a scenario that can not be relied on. An ALG can not be utilized with IPsec since there is no SIP port number in the IP payload for the ALG to key off. In addition, the inner IP header and contents are all encrypted and cannot be parsed or changed by the ALG.

undesirable delay into the communication path. In addition, both the STUN and TURN solutions require application layer support on the handset and the STUN/TURN server(s) in the service provider network must be on the public Internet, making them susceptible to threats such as denial of service attacks.

The second issue with passing VoIP traffic through a NATing gateway is that if IPsec tunneling is utilized, NATing will not be possible without the handset providing UDP Encapsulated IPsec/ESP Tunnel Mode. This is because IPsec tunneling without UDP encapsulation only presents an unencrypted IP header field, providing no UDP field for port address translation. UDP Encapsulated IPsec/ESP Tunnel Mode provides an outer IP and UDP field for NATing as shown in the following figure.

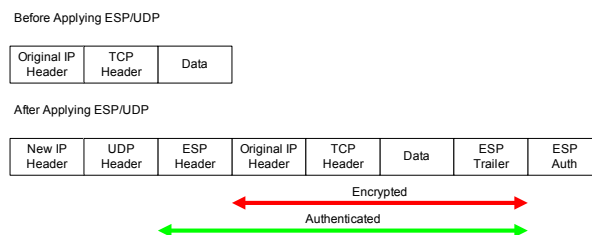


Figure 7. UDP Encapsulated IPSEC/ESP Tunnel Mode

When a NATing residential gateway receives a UDP Encapsulated IPsec/ESP tunneled packet (e.g. an RTP VoIP packet from a dual-mode handset when in a WLAN), it conducts the following operation:

- 1) replaces the Source IP Address in the outer IP packet header with the WAN IP address of the AP
- 2) replaces the Source Port in the outer UDP packet header with the UDP port selected by NAT
- 3) recomputes the IP packet Header Checksum in the outer IP packet header¹⁶

¹⁶ The Header Checksum provides a verification that the information used in processing an internet datagram

UMA and IMS require use of UDP Encapsulated IPsec/ESP Tunnel Mode for IPsec tunneled VoIP from a WLAN.

OTHER SEAMLESS MOBILITY CONCEPTS

As mentioned previously, seamless mobility in general terms is an approach that allows users to roam between application domains and communication networks without being aware of the underlying mechanisms that enable them to do so. This may include providing seamless mobility for voice and data communication from a multi-mode handset. It may also mean receiving access to a “family” number from any phone in the house or getting access to caller ID, call logs, and or billing information from a TV. In the most general sense, seamless mobility is more than seamless handover from a single device but having access to a common application or user experience from different devices at different locations. Furthermore, seamless mobility need not only apply to communications but may apply to other applications and experiences such as home monitoring and control, purchased content management, and personal content management.

In the case of home monitoring and control, seamless mobility is a means for allowing users access to their home monitoring and control system while on their home network, via a Web portal from anywhere on the Internet, from their cell phone, from their home TV/STB, and even from an automobile assistance service. The following figure illustrates these concepts.

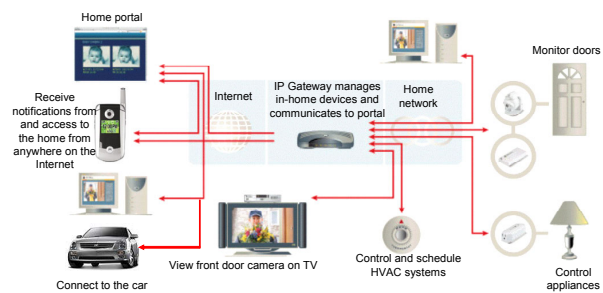


Figure 8. Home Monitoring and Control Seamless Mobility Scenario

In this example an IP gateway device in the home may be used to provide an “always on” portal between in-home devices and communications access. Home monitoring devices may have wired or wireless connections to the gateway. These devices may use non-IP protocols, in which case they are proxied on to the home network through a protocol translator to make them look like they have native IP. What is key to the concepts in this and other seamless mobility solutions is that a user’s seamless mobile experience is consistent (same “look and feel”) across the different access platforms, even with varying levels of features being accessible at different platforms.

In the case of purchased content management, seamless mobility implies a means of allowing users to purchase music and video content from home and cell phone, to store and manage content, to play content from a home stereo, TV, and cell phone, and to transfer content to CD, portable audio and video player, and automobile. The following figure illustrates these concepts.

has been transmitted correctly. The data may contain errors. If the header checksum fails, the internet datagram is discarded at once by the entity which detects the error.

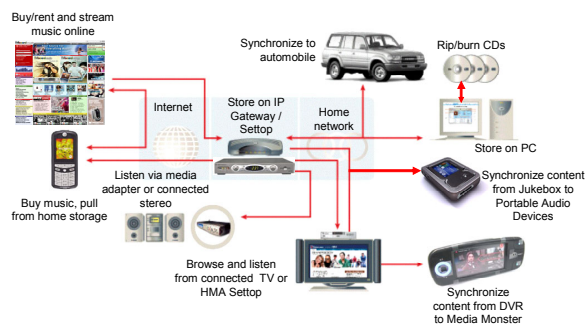


Figure 9. Purchased Content Management Seamless Mobility Scenario

In this example an IP gateway or Digital Video Recorder (DVR) Set Top Box (STB) device in the home may provide “always on” content and storage management. A partnership may be established with online content providers for seamless content purchases. Networked playback devices may be standalone or integrated. Portable audio and video players should be home network docked. Access from a gateway to automobile will need to occur via a wireless home networking connection. Digital rights management (DRM) licensing requirements may limit which devices and which network types that purchased content may have access to.

In the case of personal content management, seamless mobility implies a means of allowing users to store and manage family digital photos or video from camera, camcorder, and cell phone, to display content on TV, digital frames, or cell phone, and to order prints for digital photos over the internet. The following figure illustrates these concepts.

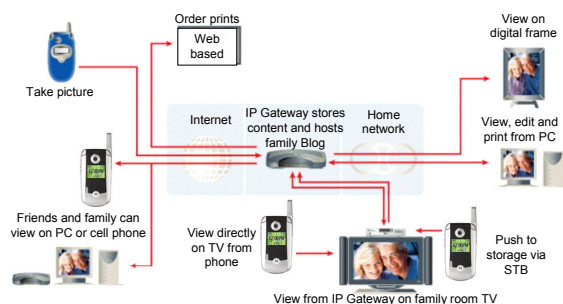


Figure 10. Personal Content Management Seamless Mobility Scenario

In this example an IP gateway device in the home may provide “always on” content and storage management as well as a family Blog. Accessory devices may be utilized for viewing photos. A partnership may also be established with online printing services for seamless print ordering.

SUMMARY AND CONCLUSIONS

Seamless mobile communications is an approach that allows users to roam between application domains and communication networks without being aware of the underlying technology or mechanisms that facilitate transparent mobility. It promises to provide value to residential and enterprise users, operators, and vendors.

Three different approaches to achieving seamless mobile communications were presented: Call Forwarding, UMA, and IMS. An IMS solution appears to be the superior choice for Cable operators.

Any seamless mobile communication solution must also account for power management, QoS, security and NAT traversal to be a complete service offering. For a dual-mode handset solution to provide acceptable battery life, additional power management techniques based on U-APSD are essential for a dual-mode handset service offering to be competitive with user expectations of cellular standby/talk times. In addition, proper handling of QoS on the WLAN and end-to-end is important in order

to provide high quality audio and differentiate service offerings between competitive VoIP service providers. Security of the WLAN and end-to-end is should be provided and WPA2-PSK and WPA2 are good solutions for residential and enterprise WLAN applications respectively. NAT traversal is also an important consideration due to the variety of NAT implementations and detection capabilities of network equipment deployed in residential and enterprise environments.

Finally, seamless mobility need not only apply to communications but may apply to other applications and experiences such as home monitoring and control, purchased content management, and personal content management.

REFERENCES

1. UMA Architecture (Stage 2) R1.0.2 (2004-11-03), Technical Specification.
2. 3GPP Group Services and System Aspects; 3GPP system to Wireless Local Area Network (WLAN) interworking; System description TS 23.234 V6.2.0 (2004-09).
3. IEEE Wireless LAN Edition, Standards Information Network IEEE Press compilation based on
4. IEEE Std 802.11TM-1999 (R2003) and its amendments.
5. IEEE Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) specifications, Amendment 6: Medium Access Control (MAC) Security Enhancements, Std 802.11i-2004.
6. IEEE Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) specifications, Amendment : Medium Access Control (MAC) Enhancements for Quality of Service (QoS) P802.11e/D12.0, November 2004.
7. IETF RFC 2474, Definition of the Differentiated Services Field (DS Field).
8. IETF RFC 3948, UDP Encapsulation of IPsec ESP Packets.
9. IETF RFC 3489, Simple Traversal of User Datagram Protocol (UDP) Through Network Address Translators (NATs).
10. IETF Internet Draft, Traversal Using Relay NAT (TURN). Work in progress.

SECURING DIGITAL CONTENT – STRENGTHS AND WEAKNESSES OF SOFTWARE AND HARDWARE IMPLEMENTATIONS

Robin Wilson
Nagravision

Abstract

Conditional Access (CA) and Digital Rights Management (DRM) are implemented in a number of ways in software (SW) and hardware (HW). Often these schemes are described as either “HW” or “SW” based security or rights management systems. Since SW requires HW on which to execute, and HW has necessarily SW running on it, the terminology is often thoroughly confusing and misleading. Since we are not dealing with locks and keys or hypothetical systems, both hardware and software elements must be present, and work together smoothly, in any sophisticated content security system.

Further confounding the confusion is the frequent use (and abuse) of terms like “replaceable”, “renewable”, “obfuscated” and “tamperproof”. In this session, these terms will be explained in the context of content security.

TERMINOLOGY AND TECHNOBABLE

Before embarking on comparing various security implementations and trade-offs, we will first define the terminology used. This step will help to overcome the ambiguous definitions and terminologies recently applied under the guise of marketing new security concepts.

HARDWARE VS. SOFTWARE SECURITY

“Overview”

The security system relies on a computation, or algorithm, to decode the protected content.

Most digital CA systems employ a unique key that enables a successful computation. The locations on the Set-Top Box (STB) for the decoding program and key are a the subject o the hardware and software security designations.

“Software Only” Security

This and other terms like “Hardware-less”, “Downloadable” and even “Renewable” are used to describe security systems where the security solution supplied by a CA or DRM company does not include hardware. Here the product of the company may be limited to only software but the inference is sometimes wrongly made that no additional resources are required or costs are incurred.

Software Needs Hardware Too

Obvious, but software still requires hardware for execution. Conventionally this hardware is referred to as the CPU. In a security system, hardware security is a significant concern.

“Hardware” Security

As we are not discussing security using tumbler locks or brass keys, it should be of no surprise that the hardware referred to here is electronic in nature and it runs software! In many instances, the decoding algorithm is programmed into the “hardware” device.

Hardware is free?

A further misrepresentation is sometimes made that a so called “SW only” systems have zero hardware cost. While it may be true that there is no hardware cost in the STB to be passed on by the CA / DRM company, someone, usually the operator, pays.

The assertion that CPU cycles are incrementally free is also misleading. Any CPU on any STB is almost always maxed out. Indeed prior to launch of a new STB or service it is unusual to find less than 110% of CPU resources are already assigned. There is a significant opportunity cost and a real cost to freeing up cycles for security applications. Even if the security application takes only 20% of an existing CPU that can become several dollars in additional cost.

As the computational power of the STB increases and the sophistication of the CA system is enhanced to meet ever-growing threats, there is the risk that legacy STB’s will be unable to support the CA system necessary to protect content.

To Summarize:

- Security software runs on hardware
- Security hardware runs software
- Hardware is always required
- Hardware can be secure or insecure
- You pay for security hardware even in a “SW” only system

Hardware can be secure or not. Likewise, security software can either be written without concern for potential reverse engineering or cleverly concealed attacks within the hardware/system. Efficient implementations are often a hybrid mix of several SW and HW security techniques.

SOFTWARE IN A SECURITY SYSTEM

Bug Free (and we really mean free)

A Bug in any software jeopardizes your hard-earned customer relationship and potentially your corporate reputation. While a bug in a security system can be catastrophic and possibly threaten your business. Even a single bug can quickly give hackers many more clues as to how a system operates compared with months or even years of analyzing a bug free system. This is further compounded in a security system where a bug can be classified as any unplanned operation regardless of the stimulus. For this reason the security of SW has much more stringent requirements including the need to do nothing when stimulated by any of an infinite array of malicious or accidental stimuli. Try requiring that of you web browser!

Beyond the impact on the CA system, any unplanned and unexpected operations will impact the viewer. This is likely to have a negative effect on the viewing experience.

Small is Good

For the reason above and for speed of operation, bug free security software is best written in very small kernels by very small teams with exhaustive regression testing.

Throwing tens or hundreds of staff at the problem will not help. Neither will bloating the code-base with huge code footprint. It is

in this area that a few very experienced programmers can easily out-perform the huge corporations who's bug ethics are driven by "good enough to ship" or "let the end customers find the bugs". It is easy to see that a large, complex, system presents many more opportunities for error and bugs.

Isolation

It is counter-productive to implement tight highly-secure bug-free security kernel only to have other applications sharing CPU or memory resources (planned or unplanned). That strategy will totally undermine security. Security will be at the mercy of the application suite du-jour. In-turn the QA issues of requiring that all the applications sharing the same code and memory space are totally bug free make this flawed implementation unworkable. Security software needs to be isolated in a protected environment isolated from a hostile memory and bugs. Denial of visibility and accessibility of the CA process is essential to protecting its secrets.

SW Obfuscation

This means that the operation or structure of the software is deliberately made non-obvious or non-intuitive to either human or machine. Although this term has recently gained some use in referring to automated transformations applied to pure software products like games, the technique has been in use for more than a decade in CA or DRM systems.

SW Tamper-proofing

Here the goal is to detect any abnormal operation in the SW due do any unwelcome external stimulus. When detected, the tampering will almost always result in either

a temporary or permanent halt to the security processes.

HARDWARE IN A SECURITY SYSTEM

Having bug free security SW is useless if the operating states, registers etc, can easily be monitored. While it is well beyond the scope of this paper to discuss the security philosophies relating to hiding and keeping secrets, having a transparent hardware platform like a generic CPU, where the operation and architecture are well understood, fatally undermines almost any security scheme.

HW Obfuscation

Just as with SW, in the context of HW the structure of silicon or functional blocks are deliberately made non-obvious or non-intuitive to either human or machine. The term camouflage may also be used. Here the silicon structure is laid out in an apparently identical manner for many of the building blocks and the critical differences are hidden deep inside an obscure silicon structure.

HW Tamper-proofing

Here the goal is to detect any abnormal probing of the silicon or functional block. Numerous techniques are employed from detection layers to radiation detectors to produce electrical anomalies.

One common measure employs fuseable links that can be burned away or destroyed after the CA programming is loaded into memory. This makes reading that code and analyzing the memory structure far more difficult.

As with SW tamper proofing, when an attempted intrusion is detected, the tampering will almost always result in either a

temporary or permanent halt to the security processes.

RENEWABILITY – WHAT DOES IT REALLY MEAN?

Downloadable

The downloadable feature within the security system authenticates or identifies a network element i.e. STB, securely communicates a downloadable solution, and launches the solution into a secured environment.

Replaceable

This has two possible meanings:

- 1). Electronic - The replaceable feature within the security system revokes the current security solution, restores the secure environment, and securely enables the downloadable feature for the replacement security solution.
- 2). Physical – Here a physical device may be replaced. Replacement is based on proper authentication, binding/paring and secure provisioning. Note: this does not always infer that the removal of the previous device. Physically replaceable hardware cuts both ways. It permits total replacement of a compromised CA system but it also permits cloning of apparently legitimate hardware solutions to receive unauthorized service.

Renewable

The renewable feature within the security system suspends the current executing version of the security solution, maintains the secure environment, and securely enables the downloadable feature the a new or upgraded version of the security solution.

Countermeasures

The countermeasures feature within the security system allows for secure and validated updates to the current executing version of the security solution. It also assists the service provider in detecting and disabling compromised platforms.

PERFORMANCE CONSIDERATIONS

Latency

Providing an easy-to-use viewing experience is critical to keeping customers happy and giving them no reason to look at other methods of content delivery into the home. Although not often thought as a factor in subscriber retention or churn it is important to ensure that viewers are never annoyed by additional channel selection delays. In the new competitive video environment, channel change delays will become a differentiating factor for service providers.

Latency in a CA system can be broadly categorized as two issues:

- 1). The first is the time taken between a subscriber's request to view a channel or view a stored file and the proper permission communicated to the security process. This could be summarized as "checking the viewing rights". In any modestly sized broadcast system, there is not enough bandwidth available to broadcast all the rights for each viewer frequently enough to avoid annoying delays of seconds if not minutes. The CA system must provide specialized configurations, storage of permissions, and communication capabilities for the timely delivery of each customer's permissions.

A system without these specialized configurations and communication capabilities relying on a two-way out of band (OOB) network such as one with a dedicated DOCISIS / DSG return path, encounters processing and round trip delays such that the system cannot scale into tens of thousands or more subscribers and guarantee to operate with the required quality of service. In addition such an architecture would have to bear the cost of the BW and support costs for the burdensome continuous OOB two way traffic. The same system limitation applies to a pure IP network. Pure IP networks are often limited in the bandwidth dedicate to an individual STB limiting the number of communications carriers available to communicate with the STB

In order to achieve satisfactory low latencies it therefore becomes necessary to selectively stream and cache the subscriber rights.

2). The second relates to the real time cryptographic time base(s) used. When a subscriber requests to view a new stream. It is considered desirable that the channel change delay from security system is well under one second it is considered desirable to be under 100mS. A figure of 200 mS is generally accepted as the delay threshold that causes irritation.

Latency in a CA system is a complex subject with operational considerations including the likelihood of network outages, installation immediacies, warehouse support etc, but it is important that the base-line operation of any large deployment is non-immediate making available resources so that installations and customer support can be given high priority.

Rights Management Matrix

This is the operational heart (as opposed to the security heart) of a CA or DRM system. It is the complex alignment and communication of the various rights as mapped onto the marketing driven needs including packaging, floating previews, Push VOD rights etc. Much of this functionality must execute on millisecond boundaries, uniquely control individual subscribers in multimillion subscriber systems yet require very low communication bandwidth and minimum latency.

It is this complex functionality, tightly linked to cryptography that is often overlooked and misunderstood, particularly in the area of new and emerging IP network where it is naively thought the routers, CMTS', DSLAM's or other edge devices can execute this function.

CA HARDWARE IMPLEMENTATION EXAMPLES

Secure Microprocessor

A secure microprocessor is a specialized CPU with numerous enhancements. The principal enhancement is a hardened hardware and software environment to safeguard against security attacks. Hardening features can include intrusion detection, camouflage cell structures, encrypted communication, and cryptographically secured memory with randomized page, address, and value construction. Today's modern secure micros can include many security specific enhancements including true random number generators, public key generation, and several security algorithm accelerators. A secure micro is an optimized platform for implementing security solutions. A secure microprocessor is typically used in a BGA,

SIM, Smartcard, MCM, USB Key and potentially CableCARD (see below).

BGA

A BGA (Ball Grid Array) is a popular type of physical IC package that provides high I/O density, small footprint, and physical security features. Historically, IC I/O has been through pins on the perimeter of the IC package, the BGA provided I/O through an array of solder ball connections on the underneath side of the IC package. Since the I/O connections are physically sandwiched between the IC and PC board, this maximizes the I/O connections, minimizes the footprint, and physically restricts access. Secure microprocessors are often implemented in a BGA package

SIM

A SIM (Subscriber Identity Module) is the security module predominantly used in GSM mobile phones. For content security applications, a SIM can be considered to be identical to a smartcard in functionality, the main difference being a smaller physical size and different insertion requirements compared to a smartcard. Because a SIM is visibly smaller than a smartcard is sometime assumed that it must be less expensive. In fact the production process is nearly identical for both, and in some instances additional steps are required to punch the SIM form factor out of a larger smartcard carrier.

Smartcard

A Smart Card (SC) is a credit card size security card containing a secure microprocessor. A SC has a wide variety of applications ranging from phone cards, digital identification devices, and standards-based satellite / cable renewable security.

There are two SC interface types: *Contact* and *Contactless*. The contact SC uses ISO-7816 standards pin connections to communicate via direct electrical contact. A contactless SC does not have contact pins but communicates via radio frequency (RF) using an embedded wire loop.

TV pass

A TvPass card is a proprietary security card from Motorola / General Instrument system providing a renewable security solution.

Proprietary solutions face cost challenges presented by non-standard manufacturing requirements and hardware with limited volume.

CableCARD

The CableCARD™ is a more substantial device, similar to a PCMCIA Type II- card designed for laptops, that slides into a slot on many newer high-end or high-definition television receivers and DVR's. The CableCARD™ eliminates the need for a cable STB (Set-Top Box) at least for the decoding function. The CableCARD™ contains a secure microprocessor, one or two-way data transceivers, and specialized circuitry to process security information and decrypt the digital content. It is, essentially, the entire CA system, hardware and software, on a removable device.

The current CableCARD being deployed has only one-way functionality. The standard for a two-way multi-stream CableCARD is in development.

MCM

A MCM (Multi-Chip Module) is structure consisting of two or more integrated circuits

interconnected within the same IC package. An MCM allows for high-density implementations with security provided by semiconductor-level integration. A typical configuration might be an audio-video decoder and descrambler packaged with a secure microprocessor. MCM technology attempts to keep critical interfaces within the chip structure away from attempts to compromise the security system

SOC

A SOC (System On Chip) is a general class of solutions allowing for high integration of many of the major subsystems within a digital STB (Set-Top Box) and likely including the security solution and the secure microprocessor. Today no realizable implementations are in current use.

General Purpose CPU

A general purpose CPU mentioned because of the compelling quantity of consumer PCs. Most if not all communications to and from the CPU are in the open and can be easily accessed. Monitoring, debug tools and expertise are widely available. Not an option to secure in a security system of high value content.

USB Key

The USB security key is used to combine a digital identity and security functions into an integrated security device. The USB

security solution spans the following features: secure digital ID passwords and digital credentials securely stored on the key and automatically presented to applications as required, authentication for third party verification including multiple factor authentication utilizes a variety of authentication methods including biometrics, one-time-passwords, digital certificates, and traditional PINs and passwords. Again a secure microprocessor can be used.

CONCLUSION

Implementing a complete security system in a STB is always a complex blend of many hardware and software techniques. Securing the STB, while critical is only one aspect of an overall security solution perhaps just the most visible “tip of the iceberg”. Developing entirely bug free SW and systems is critical for security. Throwing huge resources or groups at the problem is a recipe for interminable security compromises. Likewise utopian “SW only” systems that claim near free functionality and perfect replaceability are the technological equivalent of diet pills.

Beyond the scope of this paper, there are many additional security issues in the supply, fulfillment, and support chains that have equally challenging solutions. Remember a security system is only as good as its weakest link and cryptography is only a small but important part of a security solution.

SILICON INTEGRATION FOR NEXT GENERATION VOICE OVER CABLE CUSTOMER PREMISE EQUIPMENT

Tony Andruzzi, Bill Wallace
ARRIS
Patrick Hibbs
Product Marketing Manager, T.I.

Abstract

The voice over cable (VoCable) market is defined by both constant bit rate (CBR) and voice over internet protocol (VoIP) technologies. Operators that have previous CBR VoCable experience have provided significant input on the features required for next-generation VoIP. VoIP has gone through an initial deployment ramp in 2004 with many of the world's service providers either rolling out a commercial service or conducting extensive, large scale field trials. Feedback from this process along with industry input on the future direction of the market has provided a valuable set of features and optimizations to be included in next generation products. This paper will explore the features necessary in a next generation EMTA and the integrated, cost effective silicon necessary to support these features. The paper will explore the integration options that can reduce the component count and cost of the design. The paper will focus on factors and tradeoffs associated with silicon integration that can optimize power supply, battery charger and telephony interface costs. The paper will also explore the integration level necessary to support add on of peripheral components such as wireless LAN technology and convergence with mobile voice technology

INTRODUCTION

For the past ten years Voice over cable (VoCable) has been offered by several

multiple service providers (MSOs). VoCable deployments, starting as early as 1995, utilize traditional constant bit rate (CBR) and are the dominant method of providing cable telephony service today. More recently, deployments have started to utilize voice over internet protocol (VoIP) technology. VoCable deployments have reached nearly twelve million subscribers and less than 10% of the total subscribers are served by VoIP systems. In both CBR and VoIP systems, a key element to successful VoCable deployment is the client device installed at the service subscriber's site. This client device, referred to as customer premise equipment (CPE) has evolved over a decade of deployments. The CPE evolution has been driven by customer feedback and market trends. Valuable experience has been gained during field trials, initial service introductions and full scale deployments of VoCable CPE. This experience and feedback has provided CPE vendors valuable information to aid in the analysis and development of cost effective silicon technologies.

CPE Evolution

When first introduced, VoCable CPE devices used in CBR systems were expensive and bulky because they were developed with a Carrier Grade telco mindset, and typically deployed outdoors. As VoCable deployments grew and field experience increased, CPE vendors developed newer generations of CPE based on MSO feedback. Key factors that

needed to be addressed included cost, system availability, power consumption, and terminal equipment (ie. black phone set) compatibility.

VoIP technology offers an opportunity for the MSOs to continue to lower the cost of VoCable deployments as well as provide an open platform to add new future services. The earliest VoIP CPE were essentially DOCSIS Cablemodems with discrete voice interfaces, provided for early lab evaluations and field trials. These specialized Cablemodems are now often referred to as an embedded multimedia terminal adapter (EMTA). Many MSOs have completed lab evaluations and field trials and have begun VoIP cable telephony service introductions and plan large scale deployments of EMTAs.

VoCable CPE designed for CBR and VoIP systems consist of similar functional blocks. Both types of CPE utilize a cable MAC and PHY, a communications processor, memory, RF tuner, telephony line interface and power supply. EMTAs for VoIP systems include a voice digital signal processor (DSP) that is not required in CBR system CPE.

CPE for CBR systems typically use proprietary MAC and PHY silicon while EMTAs utilize DOCSIS cable modem MAC and PHY silicon. The communications processor, memory, RF tuner, and telephony line interface provide similar functionality for both types of CPE. The power supply circuitry is dependent on the requirements of the remaining functions.

Silicon integration has been applied to the various functional blocks of VoCable CPE for cost reduction. EMTA's have also benefited from the lower costs associated with large volume deployments of data-only Cablemodem as well. In particular, silicon vendors already offer products with the

communications processor and DOCSIS MAC and PHY functions integrated into a single ASIC. CBR systems have also enjoyed cost reduction due to silicon integration, however the pace of additional cost reduction has slowed due to the R&D investment shift to VoIP. Further opportunity still exists for additional silicon integration and lower costs for VoIP CPE.

CPE Features

Experience gained from VoCable deployments in both CBR and VoIP, along with industry input on the future direction of VoCable has helped to identify a set of key features to be considered for next generation CPE. This paper will focus on a few critical features and the cost effective silicon integrations necessary to support these features. The features to be explored are:

- low product cost
- telephony line count
- ILEC replacement
- installation environment
- mobile voice

LOW PRODUCT COST

As with most Consumer Electronics products, cost is one of the most important factors in designing a VoCable CPE. Cost is also one of the most difficult challenges facing the CPE design engineering team. Economics ultimately drives every design decision so cost tradeoffs must be considered when developing the next generation of VoCable CPE. This is especially important when considering cost reduction through silicon integration of CPE functions.

Silicon integration is not always the most effective way to optimize the CPE cost. Care must be taken to avoid adding silicon for features that may not be utilized in volume applications.

Integrated circuit (IC) packaging must also be considered. IC power density and thermal considerations limit what can affordably be integrated in a single die and a single package.

As the number of integrated features grow so does the complexity. Size and complexity determine design, test, and manufacturing of ICs, all of which impact cost.

Increased silicon density magnifies integrated circuit susceptibility to process parameter variations and noise disturbances. As the geometry of the IC decreases the designs may require more fault-tolerance which could add cost.

In the near future, complexity will increase due to a larger range of applications and higher level of design abstractions in silicon integrations. Robust solutions will be required for the integration of mixed technologies. Implications of increased integration will result in increasingly difficult cost decisions.

TELEPHONY LINE COUNT

An important feature of any VoCable CPE is the number of telephony ports that are supported. The number of telephony ports required is primarily driven by the type of service offering. Three types of services are considered:

- residential service
- multiple dwelling service
- enterprise service

Residential Line Count

Residential telephony service is by far the largest portion of the total deployment of VoCable CPE today. MSOs have been fairly

consistent in their requirement for the number of telephony lines for residential service. Early CBR system deployments typically required two telephony ports on CPE, although there are cases when up to four telephony ports are required. Four telephony lines provide a significant potential service differentiator for the MSO over the Incumbent Local Exchange Carrier (ILEC). The various MSOs deploying VoCable with CBR systems found that the actual number of telephony ports being deployed per resident ranged between 1 and 4 with the average being around 1.2 lines.

As expected, CPE requirements for VoIP systems have mirrored those of CBR systems. Two telephony lines per EMTA continue to be the volume preference. Silicon vendors have recognized this and are already taking steps to optimize their products for two telephony lines through integration. However, interest for CPE with a single telephony port is also being discussed as an option to provide a lower cost alternative to subscribers. Currently, single line CPE cost will be burdened by the silicon utilized to implement the prevalent two line CPE.

Multiple Dwelling Unit Line Count

Multiple dwelling unit (MDU) telephony service requires a higher line count than residential services. MSOs prefer to use a multi-line CPE device for economic reasons including ease of installation, footprint and equipment recovery. Modularity is also an important feature to the MSOs so that additional lines can be added easily as subscriber penetration increases.

Larger line sizes and modularity create a particular challenge for silicon integration. Cost per line is a critical factor in multi-line

CPE so silicon integration should be a consideration, but volumes may not justify the development cost. Also, it is important not to burden the cost of the residential CPE.

CPE silicon optimized to provide two telephony lines should contain “hooks” to provide the capability to easily expand the number of telephony lines supported. Telephony line expansion is possible by including the capability to support cascaded DSP ASICs. CPE silicon integration could also provide the ability to increase processor clock speeds to support multi-line applications. Continued telephony line interface integration is also under investigation.

Enterprise Line Count

Enterprise telephony service is targeted at small to medium size businesses with line counts between residential and MDU services. Telephony service can be provided to many small business subscribers with four line CPE. Modular multi-line CPE is ideal for businesses requiring more than four lines and also provides a means for expansion. A key distinction of enterprise service is a tendency toward higher telephony traffic which puts an added burden on the selection of features for silicon integration. Additional DSP resources are typically required to support frequent use of feature such as 3Way calling, voice compression, and T.38 Fax Relay.

ILEC REPLACEMENT

Most MSOs that have deployed VoCable CPE have offered facilities-based telephony services in order to compete directly with the ILEC. A key part of this strategy is to deploy VoCable CPE that provides a “primary line” telephony service. Primary line CPE requirements include:

- interfacing with a wide range of standard terminal equipment
- interfacing with existing in-house wiring
- broad range of technical specifications
- reliable battery backup with real-time status

Terminal Equipment Compatibility

Since the MSOs strategy for cable telephony includes ILEC replacement it is important that existing terminal equipment operate the same when switched to VoCable CPE. This is a critical contributor to the subscriber’s experience. Much has been learned regarding terminal equipment compatibility from the extensive deployments of CPE in CBR systems. Early CBR deployments were affected by a few terminal equipment incompatibilities, such as:

- handset volume and side-tone
- fax machine issues
- caller ID issues
- analog modem issues

Early CPE telephony line interfaces for CBR systems were designed to meet the Bellcore TR-909 specification. TR-909 specifies different values for some of the voice and signaling parameters than those typically seen on traditional ILEC lines. TR-909 is a specification targeting short subscriber loops like those expected in cable telephony applications. Unfortunately, some terminal equipment is impacted by these parametric differences and the equipment either worked unsatisfactorily or not at all. In CBR systems these issues have, for the most part, been resolved.

It is not surprising that some of these same issues are appearing in early

deployments of VoIP cable telephony, and in most cases for completely different reasons. This is primarily due to the fact that the CPE for VoIP systems is required to perform many of the audio functions that were handled by the public switched telephone network (PSTN) in CBR systems. These audio functions include:

- call progress signaling
- tone generation
- tone detection
- voice activity detection
- echo cancellation
- voice compression
- fax relay

Since VoIP is based on packet switched technology, a DSP is required in EMTA designs to perform these audio functions. The DSP used in the MTA must be carefully selected to insure that the required audio features can be supported on all EMTA telephony lines simultaneously. This is particularly true for two and four line residential EMTAs. In multi-line applications it may be possible to allow for a traffic model that requires simultaneous support for a percentage of the total number of available lines. This is more likely the case for multiple dwelling EMTAs but less true for EMTAs expected to support enterprise services.

Optimizing a DSP for integration with other EMTA functional blocks becomes a challenge when the audio functions are considered. VoIP CPE deployment forecasts indicate that an optimal DSP would support the full audio feature set for two telephony lines.

Particular audio functions must be considered when determining DSP requirements for silicon integration. The number of simultaneous active telephony

lines and call features supported dictate the voice codec and compression bandwidth requirements. If three-way calling is a required call feature then the voice codec processing requirements are doubled. Low bit rate (LBR) compression also adds to this processing requirement. The number of simultaneous telephony lines requiring LBR compression must also be considered.

The expected EMTA volumes support the economic integration of the DSP with the cable modem ASIC, which consists of a communications processor and a DOCSIS MAC and PHY. The size of DSP and supporting memory depends on the audio function optimization for the number of lines and call features requirements. DSP requirements are also impacted by the amount of margin in performance desired. Field upgrades for new or modified call features are possible if the DSP is sized appropriately.

The telephone line interface is also instrumental in providing terminal equipment compatibility. The line interface circuits provide the BORSCHT functions:

- **B**attery feed
- **O**ver-voltage
- **R**inging
- **S**upervision
- **C**odec
- **H**ybrid
- **T**est

The BORSCHT functions are implemented primarily with a voice signal processor and a subscriber line interface circuit (SLIC). The line interface must be software configurable to adequately address terminal equipment compatibility. The voice signaling processor provides channel filtering, input impedance synthesis, transhybrid balancing, gain adjustment, and

voice path diagnostics. The SLIC provides the high voltage interface for battery feed and ringing, as well as the two-to-four wire conversion and loop test.

Various options exist for telephony line interface silicon integration. The voice signal processor is often implemented with a DSP so that a suitable integration with the DSP is possible, thereby providing the audio features described earlier. The difficulty is determining if this is a cost effective integration. Expanding the capabilities of a DSP to provide both the audio features and the voice signaling may not provide a cost advantage over discrete components. Integration of DSP functions can be a cost advantage if the DSP is properly optimized and volumes are high enough. Concerns with integrating the DSP functions are possible feature reduction due to memory size constraints and the number of telephony lines supported.

Alternately, the voice signaling processor could be integrated with the SLIC. This option has the disadvantage of mixing low and high voltage silicon technologies but is ideal for multi-line telephony applications.

Service Availability

An important requirement of primary line service is to provide telephony service during utility power outages. In CBR systems this was implemented with three distinct CPE powering schemes:

- network power
- local external uninterruptible power supplies (UPS)
- internal battery backup

For today's CPE, the power system design and the use of low power, high performance subsystems has the greatest effect on achieving the powering goals of primary line

service. Since the CPE power consumption directly affects battery backup time, a highly efficient, low power design is a must. The subsystems, including the cable modem IC, battery charger, telephony interface and RF tuner, must continue to provide high performance while consuming as little power as possible.

Many of the power system components can be readily integrated into the cable modem IC for improved cost. To make this successful, attention must be directed to low noise designs, versatile circuit blocks, and high power efficiency.

Linear regulation is perhaps the easiest function to integrate and is available in today's cable modem ASIC's. These regulators should operate with low dropout voltage requirements to maximize efficiency and provide excellent ripple and noise rejection across a wide range of frequencies. The output voltage of these regulators should be available outside of the cable modem ASIC for powering external circuits. The available current from these regulators should allow powering of typical CPE functions. Thermal limitations should also be addressed within the ASIC.

The efficiency demands of certain voltage rails will preclude the use of a linear regulator. For these requirements, the integration of a generic Pulse Width Modulator (PWM) function will allow the power systems designer to design a highly efficient (>90%) switching converter using external power components. This PWM function must be versatile to allow implementation flexibility for the system designer. Within the cable modem IC, care must be taken to minimize noise effects from the PWM and associated external power components on the cable modem functions.

The battery charger and monitoring circuits are excellent candidates for silicon integration. Today's chargers are implemented with simple 8-bit microprocessors, A/D converters, Op amps and similar circuit functions. These items are inherently simple to integrate into the silicon and are low power by design. The implementation within the silicon must maintain flexibility to allow charging of multiple battery chemistries, voltages, currents and algorithms.

Finally, silicon integration should take advantage of power management features. The integrated components can be controlled to provide the lowest power consumption during various states of operation. During idle modes, selected IC functions could be shut down providing maximum battery back up efficiency. Processor clock speed should also be controllable to allow optimal power consumption. The processor clock could then be slowed during idle modes to save power.

INSTALLATION ENVIRONMENT

VoCable CPE deployed in CBR cable telephony systems are principally installed outdoors, either on the side of a building, on a pole, or occasionally strand-mounted. MSOs preferred this outdoor method to simplify installation and maintenance. Outdoor CPE are required if network power is the preferred powering scheme.

VoCable deployments have shifted towards indoor EMTA installation to take advantage of lower product, but outdoor EMTAs are being considered for some residential and multi-line deployments. ASIC operating temperature becomes an important aspect of silicon integration if outdoor deployments are supported. Silicon cost is impacted by extended temperature

requirements due to compensation circuits and expanded test times. There is also the possibility that certain integrations are not cost effective at extended temperatures.

MOBILE VOICE

Several additional features related to mobile voice services are worth examining related to EMTA silicon integration. These features include implementing cellular codecs in the audio feature DSP and adding wireless capabilities to the EMTA.

Each time a voice packet is translated from one codec format to a second codec format, voice quality is degraded. This codec format translation is referred to as transcoding. A call between a cell phone and an EMTA goes through at least two transcodings. For example, the call would be translated from GSM to G729 and then from G7.29 to GSM. As the number of transitions increases the voice quality decreases.

Instead of compressing voice traffic using the standard set of cable-centric codecs that are transcoded at network gateways, an EMTA could use a cellular based codec. This would eliminate a transcoder operation that contributes to lower voice quality.

Codecs are audio functions so they are implemented within the DSP. The addition of cellular based or any other native equipment codecs impact the DSP requirements and must be considered when planning silicon integration.

At some point market economics will support the integration of silicon for wireless local area network (WLAN) applications. Broadband-enabled phones will support high-quality audio, live video streaming, full-speed web access, automatic and instantaneous synchronization with address

books and email services, and real-time internet gaming.

An EMTA's ability to support voice-over-WLAN also enables convergence of the cable and cellular networks. Subscribers will use the traditional cellular network when they are mobile, and access the VoIP network when they are in the office or at home, utilizing the same dual-mode cellular/Wi-Fi handset.

SUMMARY

Much has been learned during the past decade of VoCable deployments, and there is still more to learn with the advent of VoIP cable telephony. What has been learned can readily be applied to designs for the next generation VoCable CPE. The next generation EMTA designs must take advantage of silicon integration opportunities to meet the cost requirements of the MSOs. Many options exist for integrating EMTA

functions into single ASICs cost effectively. The key to a successful cost reduction is to carefully select the features, scaling capabilities, upgrades and future expansions. Timing is also critical in cost effective silicon integration. It is important to apply the proper level of integration for optimal performance for a given market.

REFERENCES

1. M. Paxton, "Cable Telephony Service: VoIP Finally Shows Up", In-StatMDR Report No. IN0401244MB, December, 2004.
2. W. Maly, "Cost of Silicon Viewed from VLSI Design Perspective, Electrical and Computer Engineering, Carnegie Mellon University.
3. W. Maly, "A Point of View on the Future of IC Design, Testing and Manufacturing", Electrical and Computer Engineering, Carnegie Mellon University.

SOME CONSIDERATIONS IN THE USE OF FORWARD ERROR CORRECTION (FEC) IN BROADBAND VIDEO DISTRIBUTION OVER WIRELESS LANS

Walter Boyles
Consultant

Abstract

MSOs seeking to deploy broadband video services to subscribers who are using wireless LANs should consider the use of forward error correction (FEC) technologies as well as the pros and cons of the use of streamed video versus file-based video delivery.

INTRODUCTION

As MSOs look to deploy new services to broadband subscribers, they find a number of home networking products and technologies in place in their subscribers homes in much the same way that MSOs find a variety of PCs with different CPUs, memory configurations, and operating systems. The MSO is likely to seek to deploy services such as broadband video to as much of the equipment in place as is reasonably possible to accommodate their subscribers.

802.11

The majority of WLAN products in home use employ one form of the 802.11 standards. The more well known versions of 802.11 are summarized in Table 1.

Table 1 - Summary of Some 802.11 Standards

802.11 Standard	Modulation	Freq Band	Max Link Rate	Theo Max TCP Rate	Theo Max UDP Rate
802.11b	CCK	2.4 GHz	11 Mbps	5.9 Mbps	71. Mbps
802.11g (w/ 802.11b)	OFDM/ CCK	2.4 GHz	54 Mbps	14.4 Mbps	19.5 Mbps
802.11g (11g only)	OFDM/ CCK	2.4 GHz	54 Mbps	24.4 Mbps	30.5 Mbps
802.11a	OFDM	5.2, 5.8 GHz	54 Mbps	24.4 Mbps	30.5 Mbps

The actual performance of the wireless LAN, of course, will be complicated by a huge array of factors including (but not limited to) the distance within the home, the configuration of the room, the number of walls and other solid objects between the access points and the client, and even RF interference. The 2.4 GHz modes of 802.11 can incur interference from a variety of other appliances in the home including cordless phones and microwave ovens.

In fact, even other factors such as packet size will affect the performance of the wireless network, as well.¹ Of course, the networking equipment itself which relies on chip sets in different versions from a variety of vendors will be one of the largest determining factors in the performance of the wireless network.

In open office environments, the observed performance of 802.11a using equipment

based on previous generation semiconductor technology generally maintained at least 1 Mbps for distances of up to almost 60 feet and approximately 3 times that distance for equipment using a newer version of the chipset.²

The performance of the more common 802.11b in a similar open office environment generally stayed above 1 Mbps for up to about 140 feet.² However in general usage such as in enclosed offices with walls, closed doors, and other solid objects common 802.11b home networking equipment, commercially available in 2003 and 2004, was not reliable at data rates above 700 kb/s.

Forward Error Correction:

The idea of using forward error correction (FEC) in communications is of course well-established and dates back to the birth of information theory and its legendary inventor Robert Shannon. Shannon, of course, proved that there is a capacity for any communications channel and that reliable communications is possible for rates approaching that theoretical capacity. That theoretical capacity is determined as follows:

$$C = W \log_2 (1 + P/N)$$

where:

C is the channel capacity in (in bits/sec)

W is the bandwidth (in Hertz)

P is the transmitter power (in Watts)

N is the noise (Watts)

From information theory, techniques emerge to model a channel that can be used to predict the probability of an error for N bytes of packet length such as below:

$$P_e^m(N) = (1 - p_b^m)^{8N}$$

where p_b^m is the bit error rate (BER) of the PHY mode at a given SNR (signal to noise ratio).

Bypassing most of the mathematics, in this discussion Shannon introduced the idea of codes as ensembles of vectors that are to be transmitted. If the number of vectors is $K=2^k$ each vector can be described with k bits. The vectors are assumed to be of equal length and we refer to this length as the block length. For a length of vectors “n”, then n times k bits has been transmitted and the resulting code of has a rate of k/n bits per channel.³

Shannon went on to prove the existence of these codes that allow us to approach the capacity of the channel and since his groundbreaking work, a variety of coding techniques have been developed. This ended the concept that reliable communications meant that transmission power must be increased and/or messages must be sent repeatedly. By the early 1990s, the best coding solutions had achieved actual capacities within 3-5 dB of the theoretical limit shown by Shannon.

A breakthrough occurred in 1993, when two French engineers, Berrou and Glavieux, introduced turbo coding that resulted in a 3dB improvement over existing coding schemes. Turbo codes can be classified as turbo convolutional codes (TCC) and turbo product codes (TPCs).³

As coding technology advanced and the desire for better and simpler codes increased a class of codes previously discovered by Robert Gallager some 30 years earlier reemerged. This coding scheme called LDPC – or low-density parity checking codes perform closer still to the Shannon limit.

A comparison of LDPC code to PCCCs (parallel concatenated convolutional codes), SCCC (serially concatenated convolutional codes), and TPCs is shown in Table 2.

Table 2 - Comparison of Turbo Coding Schemes⁴

Attribute	PCCC	SCCC	LDPC	TPC
Code Rate	fair	poor	good	poor
Block Size	good	good	poor	poor
Modulation	good	best	good	good
Complexity (throughput)	fair	fair	good	good
Performance	poor/ best	good/ fair	fair/ good	fair/ good

A large number of researchers worked on LDPC codes which has also resulted in a variety of vendors for coding products that use LDPC. Some of these additions to the basic LDPC theory bear the names of the researcher or company that has developed them and each puts forth some advantages over the base LDPC coding developed by Gallager. However, LDPC coding (as shown in Table 2) provides a good compromise of attributes for use in file-based broadband video distribution over wireless LANs.

The codes can be implemented to utilize variable amounts of overhead to provide different degrees of error correction depending on the noise on the channel – overhead may typically vary from 5%-50% and provide resistance to data loss.

Considerations of Actual Use of FEC:

Now, in actual use, the success of achieving data integrity with a coding scheme such as LDPC is due to a number of factors (some mentioned here) as well as the overhead and the particular coding scheme chosen. Of course, one critical issue is the way the video content itself is distributed.

LDPC coding works particularly well for file distribution. There is an added advantage in the distribution of large files since the use of a large file mean your code is operating over a large number of bits. Thus, the loss of any small number of bits is likely to be corrected quite easily with an extremely high probability.

Take the example of a one hundred (100) packet message, that is transmitted, using an LDPC code with 10% overhead. The loss of say any 7 packets has less than a 10^{-6} chance of resulting making the video file unviewable. Or put another way, the probability that any particular 7 message packets that are lost would result in a correction that could not be performed is less than 1 in a million.⁵

Now, of course, there are other issues with regard to packet loss. One view is that packets lost will typically be the same from one subscriber as another in an IP Multicast distribution scenario. So, the idea that you can request replacement packets that were lost and correct the files that way has been suggested.

In this author's view, while there are undoubtedly network incidents that occur which result in the loss of the same packet or packet for some number of subscribers, there are also random events not only at different points in the network but on each subscribers computer, as well. These random and singular losses are better corrected by using an LDPC coding scheme with say 5%-10% overhead. This amount of overhead is barely noticeable and in the case of IP Multicast distribution the overall efficiency of content distribution is so higher anyway as opposed to unicast file distribution that an additional 10% of overhead is a very small price to pay.

Of course, FEC can also be used in video streaming. The one penalty incurred in using FEC for streamed video is that a delay is imposed on the stream to perform the error correction calculations over some number of packets. The issues regarding what is an acceptable delay are dependent upon a number of factors including amount of delay is tolerable to the viewers.

Then again, in a streamed video environment, the loss of a few packets may not be significant in the same way as it is in file-based video delivery. The loss of few packets might only momentarily degrade the signal and still make the quality of the video experience acceptable to the subscriber.

Again, the question of what is acceptable is dependent upon the many factors from the video quality that is trying to be achieved to the wireless LAN performance from the access point to that particular subscriber to the length (in time) of the disruption, as well as what packets are lost and how frequently the loss reoccurs.

The delay considerations may also make the preferred coding scheme in streamed video different than that of the coding scheme in file-based video distribution.

One such proposal for video streaming is the system for packet loss protection for the H26L-FGS. This employs Reed Solomon coding along with a system of requests and acknowledgements. The combination of feedback and error correction may be more suitable than in streaming than the FEC coding schemes used in file-based video delivery.⁶

CONCLUSION:

Overall, however, the one factor that can not be overcome in a wireless LAN is of course the network performance. If you can not reliably distribute data at over 700 Kb/s, due to whatever factors, then the use of FEC is not going to allow you to distribute video at a higher data rate whether it is streamed or delivered file-based.

For streamed video, the limitations of the wireless LAN performance will of course also limit the encoding rate you can choose for streamed video. You will always be limited in video resolution by the lowest common denominator of what data rates the network can tolerate.

A file-based video distribution system has no such limitations. In fact, even across a network with a 700 kb/s or less bandwidth limitation, it is easily possible to distribute files that are encoded at high definition quality. The use of FEC and IP Multicast makes such distribution across home wireless networks in place today entirely possible and implementable.

Table 3. summarizes some of the considerations in video delivery methods and the use of forward error correction of current wireless LANs. From the table, it can be seen that the combination of FEC with file-based delivery and is very efficient – especially when these files are delivered with IP Multicast.

Table 3 - Comparison of Video Delivery and Quality w/ FEC.

Video Distribution Method	Tolerance to Minimal Packet Loss	FEC Delay	FEC Value to Video Quality	Achievable Video Quality over Most WLANs
Streamed w/out FEC	Fair	None	N/A	Worst
File-Based w/out FEC	Poor	None	N/A	Fair
Streamed w/ FEC	Good	Worst	Fair	Fair/poor
File-Based W/ FEC	Best	Not Noticed	Best	Best

The author can be contacted at walterboyles@comcast.net.

Bibliography:

1. Overhead Constrained Packet Video, Multimedia Communications Laboratory, University of Texas at Dallas, 2002.
2. 802.11 Wireless LAN Performance, Atheros Communications, 2003.
3. LDPC: Another Step Toward Shannon, Tony Summers, CommsDesign.com (EE Times), Oct., 2004.
4. Forward Error Correction, TrellisWare Technologies, 2005.
5. Private Conversation, Bhavan Shah, 2005.
6. Packet Loss Protection of Scalable Video Bitstreams Using Forward Error Correction and Feedback, Charfi and Hamazoui, 2003.

SUB-BAND DIVISION MULTIPLEXING (SDM) INCREASES BANDWIDTH EFFICIENCY AND PROVIDES HIGHER TOLERANCE OF COMPOSITE DISTORTIONS

Mark E. Laubach, Yi Ling, William J. Miller and Tracy R. Hall
Broadband Physics Inc.

ABSTRACT

Bandwidth needs of customers on cable plants have increased dramatically over recent years and will continue to rise in the near future. Increasing the order of QAM modulation has been the most popular way to satisfying the bandwidth needs until recently when nonlinear distortions and limited dynamic range in HFC systems have proved to be an obstacle for reliable 256 QAM service. Hence, it is of both practically necessary and theoretically interesting to investigate approaches other than QAM to increase bandwidth efficiency and to provide higher tolerance of composite distortions all at the same time. Sub-band Division Multiplexing (SDM) is one of these new approaches. In this paper, we give an introduction of the Sub-band Division Multiplexing (SDM) technique based on filter bank scheme and wavelet mathematics. SDM represents a philosophical change comparing to standard QAM in terms of baseband signal formulation and alphabets selection. To show this change, the fundamentals of SDM will be overviewed. Unique characteristics resulting from the SDM fundamentals will also be presented. It will be shown that these characteristics implying multiple advantages of SDM over equivalent QAM on cable applications, especially the tolerance of composite distortions. Simulation results and measurements on the Broadband Physics prototype system from the lab and field trials also will be presented in the paper to verify the theoretical results.

I. INTRODUCTION

SDM stands for Sub-band Division Multiplexing. It is a technique of dividing RF spectrum into multiple and equal-sized sub-bands using filter bank structures. When the basic sub-band filter is designed to have steep roll-off property, the spectrum of each sub-band has little overlap into the neighboring sub-bands. Combined with SDM's inherent frequency and time orthogonal basis, each sub-band is highly independent of other sub-bands, in both frequency and time. For this independency property, SDM can be used as a digital signaling scheme to transmit data stream in parallel over the multiple sub-bands that it creates without having large inter sub-band interference. Used as such, SDM is a multi-band carrier-less modulation. Each sub-band operates in the same manner, but with a different offset frequency.

SDM can be applied to digital communication over many different types of media, such as cable [1], power line and wireless. The first proposal of applying SDM in high-speed data communication was made almost ten years ago by Miller [2], the founder of Broadband Physics Inc. Later, similar ideas of using filter bank techniques for data communication were proposed in other literatures, such as [3]. However, since the 1990s, Broadband Physics Inc. has been the active leader in developing the SDM technology to be implemented in a variety of applications

over different channels. Currently, Broadband Physics Inc. is focusing on developing SDM modems for cable downstream applications. It will be argued in the sequel that SDM has numerous benefits in the cable downstream applications. The most prominent two are increased bandwidth efficiency and higher tolerance of composite distortions. In its implemented form, SDM also demonstrates many advantages over popular digital modulation schemes other than QAM, such as Orthogonal Frequency Division Multiplexing (OFDM). These advantages include but not limited to achievable higher bandwidth efficiency with lower system complexity, less sensitive to phase noise, highly resilient to multi-path impairment. All these advantages make the SDM particularly well suited for wireless applications as well. However, this paper discusses all the said characteristics of the SDM but focuses on its high bandwidth efficiency and tolerance of composite distortions for cable applications. The rest of the paper is organized as the following:

Section II gives an overview of the fundamentals of the SDM and discusses some unique characteristics resulting from those fundamentals. This section further discusses the advantages of the SDM closely connected with its characteristics.

Section III discusses the bandwidth efficiency of SDM and through an example, shows that SDM has higher actual bandwidth efficiency than the theoretically equivalent QAM.

Section IV discusses the resiliency of SDM to composite distortions in depth.

Section V presents the results from a Broadband Physics Inc. prototype system test in a simulated cable plant.

Section VI concludes the paper.

II. SDM AND ITS CHARACTERISTICS

Based on a well-designed band-pass filter with high stop-band attenuation, a filter bank can be constructed by frequency shifting this filter prototype and combining them as a polyphase filter. Assuming that the transfer function of a prototype band pass filter is

$$F_0(z) = \sum_{n=-L}^L h_0(n)z^{-n} \quad (1)$$

where $z = e^{sT}$ and T is the sampling time, if we over-sample it at M times the original sampling rate, then the over-sampled version of $F_0(z)$ is

$$F(z) = \sum_{k=-ML}^{ML+M-1} h(k)z^{-k} \quad (2)$$

where $z = e^{s \frac{T}{M}}$.

For $k = nM, n = -L, \dots, 0, \dots, L$, we have

$$h(k) = h(nM) = h_0(n). \quad (3)$$

Separating $F(z)$ into its polyphase components, we obtain

$$\begin{aligned} F(z) &= \sum_{k=-L}^L (h(kM)z^{-kM} + h(kM+1)z^{-(kM+1)} \\ &\quad + \dots \\ &\quad + h(kM+M-1)z^{-(kM+M-1)}) \\ &= F_0(z) + z^{-1}F_1(z) + \dots \\ &\quad + z^{-(M-1)}F_{M-1}(z), \end{aligned} \quad (4)$$

where $F_0(z)$ is the same as defined in (1) because of (3), and its polyphase versions are

$$\begin{aligned} F_1(z) &= \sum_{k=-L}^L h(kM+1)z^{-kM}, \\ F_2(z) &= \sum_{k=-L}^L h(kM+2)z^{-kM}, \dots, \\ F_{M-1}(z) &= \sum_{k=-L}^L h(kM+M-1)z^{-kM}, \end{aligned} \quad (5)$$

for $z = e^{j\frac{T}{M}\omega}$.

Here we can construct a filter bank using the prototype $F_0(z)$ and its frequency-shifted versions of $F_1(z), F_2(z), \dots, F_{M-1}(z)$ by choosing the prototype so its frequency shifted versions form an M-band quadrature-mirror filter bank, the impulse responses for each sub-band filter are wavelets orthogonal in both time and frequency. The polyphase construction makes it computationally efficient. With this filter bank, we can transmit data in parallel through all or part of the M branches of the filter bank with each branch being delayed by one sample clock $\frac{T}{M}$ from the previous branch. Due to the clock delay between each branch, we can combine the outputs of each branch to form a single transmitting signal. To express this combined transmitting signal in equation, we assume that branches

$$F_j(z), \dots, F_{j+l-1}(z), 0 \leq j, \dots, j+l-1 \leq M-1,$$

are used to transmit data. The data streams with rate $\frac{1}{T}$ for these branches are defined as

$$\begin{aligned} A_j &: a_{j0}, a_{j1}, \dots \\ A_{j+1} &: a_{(j+1)0}, \dots \\ &\dots \\ A_{j+l-1} &: a_{(j+l-1)0}, \dots \end{aligned} \quad (6)$$

These data streams can come from different alphabets or constellations. For instance, A_j comes from a four state Amplitude Modulation (4-AM) alphabet, taking values from $\{-3, -1, 1, 3\}$, while A_{j+1} comes from 8-AM, taking values from $\{-7, -5, -3, -1, 1, 3, 5, 7\}$, etc. Delaying A_{j+d} with respect to A_{j+d-1} by $\frac{T}{M}$ for $0 \leq d \leq l$, and combining with all zero data streams for the unused branches, we have the following combined data stream with rate $\frac{M}{T}$

$$\begin{aligned} A &: 0, \dots, 0, a_{j0}, a_{(j+1)0}, \dots, a_{(j+l-1)0}, 0, \dots, 0, \\ &0, \dots, 0, a_{j1}, \dots, a_{(j+l-1)1}, 0, \dots, 0, \\ &0, \dots, 0, \dots \end{aligned} \quad (7)$$

Transmitting this combined data stream through the complete filter bank $F(z)$ is equivalent to transmitting $A_j, \dots, A_{j+l-1}$ (expanded to rate $\frac{M}{T}$ through expanders) separately through the corresponding branches, delaying by one clock from each other and combining them at the output as shown by the equation below

$$F(z)[A] = \sum_{k=j}^{j+l-1} z^{-k} F_k(z)[A_{ke}] \quad (8)$$

where the square bracket denotes the filtering operation and $A_{ke}, j \leq k \leq j+l-1$ are the expanded $A_j, \dots, A_{j+l-1}$.

We can also illustrate this process in spectrum plots.

Figure 1-1 shows an example spectrum of the output of a single branch.

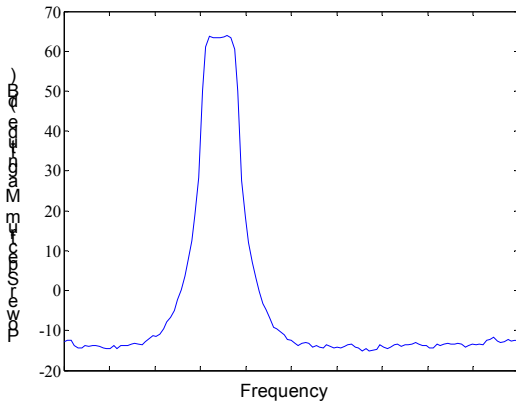


Figure 1-1: Single sub-band spectrum

Figure 1-2 shows an example spectrum of the combined outputs from all active branches.

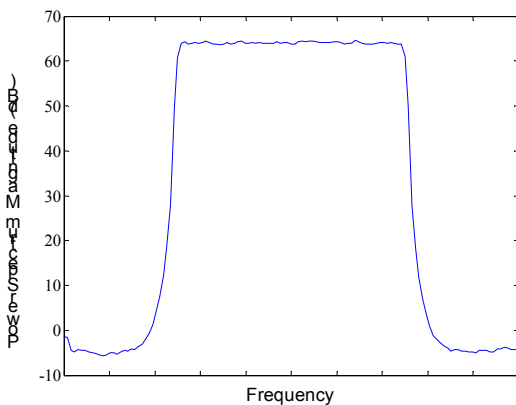


Figure 1-2: Combined multi sub-band spectrum

In view of the spectrum plots in Figure 1-1 and 1-2, we see why the name “Sub-band Division Multiplexing”, each branch creates a sub-band of the entire spectrum; the overall spectrum is the combination of all the sub-bands. The performance shown in Figure 1-2 has been achieved in digital hardware by Broadband Physics, Inc. With such transmitter structure, we can easily see that the receiver structure is essentially the same as the transmitter with the same filter bank structure and correspondingly matched filters filtering the received signal into separate sub-bands. The

orthogonality of the filters can be compromised if, for example, there is a group delay variation across the subbands. The resulting self-interferences are relatively easy to remove if the prototype filter has very steep roll-offs. Another useful view of the above transmitter and receiver structure is the wavelet transform. The transmitting filter bank is a form of inverse wavelet transform.. Each data symbol is amplitude modulating a wavelet, the data stream to be transmitted is a signal vector comprised of modulated wavelets or equivalently, a vector expressed by a set of base functions in the wavelet transform domain. Hence, the receiving filter bank just needs to be a wavelet transform [2], [4].

Comparing to other typical digital modulation schemes, such as QAM, SDM has several unique characteristics. All the characteristics of the SDM are the results of the basic construction of the SDM as shown above. Furthermore, these characteristics are the reasons behind the advantages of SDM comparing to other typical digital modulation schemes. Here we discuss some of them. For more detailed comparisons between SDM and other modulation schemes, e.g. QAM and OFDM, please see [5].

First, depending on the design of the prototype filter, 50dB plus stop-band attenuation is achievable by the SDM spectrum. This characteristic of the SDM makes an extra pulse-shaping filter unnecessary at the transmitter. With typical additional 10% bandwidth required for a pulse-shaping filter, (e.g. raised cosine filter) for QAM, SDM with the same theoretical bandwidth efficiency has more effective bandwidth efficiency.

Second, each data symbol is confined to its respective sub-band by the filter bank as

seen in Figure 1; consequently, the abrupt transitions of the data symbols do not cause ringing in the channel unlike the multi-carriers in OFDM. Since no ringing effect exists, cyclic prefix is not needed in SDM. In OFDM the length of the cyclic prefix can sometimes exceed 30% of the symbol duration. Thus SDM can achieve significantly higher effective bandwidth efficiency comparing to OFDM. Like OFDM, the bandwidth of each SDM sub-band and the number of sub-bands are design parameters and can be optimized accordingly for the channel. Furthermore, by selectively choosing active and inactive sub-bands, we can tailor the entire transmitting spectrum to the actual operating channel. The benefits of such flexibility are numerous, including easy fitting under emission masks and easy mitigating narrow band interference [4]. Because of the very high stopband attenuation of SDM subbands there is an advantage compared to OFDM, which typically has only 13 dB attenuation between adjacent bands.

Third, SDM can transmit a real signal, i.e. signal with only I component but no Q component without being limited to lower bandwidth efficiency. This can be achieved by choosing the data stream for each sub-band from some amplitude modulation (AM) constellations (same constellation is not required for different sub-band). Since the transmitted signal is a real signal, the decision region of the receiving slicer is single dimension, it can stand up to higher phase noise than signals with both I and Q components carrying the information. Fourth, SDM is orthogonal in time allowing the overlap in the transmission of one symbol with previously transmitted symbols [5]. The amount of symbol to symbol overlap is a design choice.

Fifth, recalling that the SDM transmitter and receiver pair is essentially a wavelet transform pair. Any equalization to be done in the receiver while following the receiving filter bank is not in time domain. Thus the equalizer structure can be made significantly less complicated than the time domain adaptive equalizer typically seen in QAM systems.

The above five properties and advantages of SDM are not the only benefits of using SDM, rather they are the five most essential and also most intuitive to describe at this time.

A note of notation is helpful here and for the rest of the paper, depending on the number of bits a symbol in each subband represents, the corresponding SDM scheme is called L -SDM for L bits per symbol in a single subband, or equivalently, for an alphabet of size $S = 2^L$ for the subband.

In the next two sections we will discuss advantages of SDM in cable applications, especially its capability of providing actual higher bandwidth efficiency and mitigating composite distortions.

III. SDM PROVIDES INCREASED ACTUAL BANDWIDTH EFFICIENCY

Since one of the current major driving forces of digital cable technology development is the need of bandwidth, any modulation scheme that can provide higher than current standard bandwidth efficiency (measured in bits/second per Hz or bps/Hz) will increase the raw digital capacity of the spectrum without increasing the actual bandwidth. SDM happens to be such a modulation scheme.

To study the bandwidth efficiency of the SDM, we can start with an example for cable applications. Each current cable

channel is 6 MHz wide. To formulate a baseband transmitting signal on one channel using SDM, we can first choose the digital sampling frequency to be $2f$ ($f > 6$ MHz to satisfy the Nyquist theorem). Then, we divide the spectrum of $0-f$ Hz into M subbands, each subband has bandwidth of f/M and the symbol rate for each subband is $2f/M$. For convenience, f can be chosen such that f/M is an integer. We can turn on any block of continuous $6 \times 10^6 \times M/f$ subband to create a 6 MHz cable channel. Two of the active subbands at both edge of the channel can be turned off to avoid interfere with the neighboring channels.

Again, for the convenience of presentation, we assume that each active subband has the same alphabet that has S states. So the bandwidth efficiency of each subband is $\log_2(S) \times 2f/M / (f/M) = 2 \times \log_2(S)$ bps/Hz, therefore, the bandwidth efficiency of the combined channel is also $2 \times \log_2(S)$ bps/Hz if we just ignore the two inactive guard subbands at the channel edges for now.

With the general derivations above, we can look at some actual numbers for the example. Assume that the chosen digital sampling frequency is 51.2 MHz, the entire useable bandwidth for SDM is $51.2/2 = 25.6$ MHz. By using SDM technique, we can divide 25.6 MHz into 256 subbands; each subband is 100 kHz wide with a symbol rate of 200 kHz. To create a 6 MHz wide cable channel, we can choose any 60 continuous subbands out of the entire 256. Please note that the entire 256 subbands are for signal formulation purpose only, no actual energy being put on them except the 60 chosen subbands, the real occupied spectrum is still only 6 MHz wide, the other unoccupied spectrum is free for other uses. Suppose we choose 60 subbands from 9 MHz to 15 MHz, two edge subbands being turned off as

guard bands, so the actual number of subbands for data transmission is $60 - 2 = 58$. For the convenience of presentation here, assume that each subband uses the same 16-AM alphabet with 16 states on a single real axis. Each subband has bandwidth efficiency of $\log_2(16) \times 2 \times 100,000 / 100,000 = \log_2(16) \times 2 = 8$ bps/Hz, the same as 256 QAM. The actual overall combined channel data rate (accounting the two inactive guard subbands) is $58 \times 8 \times 100,000 = 46.4$ Mbps, the actual overall channel bandwidth efficiency is $58 \times 8 \times 100,000 / (60 \times 100,000) = 58/60 \times 8 = 7.7$ bps/Hz, which is less than 4% lower than the theoretical bandwidth efficiency of 256 QAM. However, considering the excess bandwidth of QAM from the pulse shaping filters such as root raised cosine filter, the actual 256 QAM bandwidth efficiency can be 10% less than the theoretical value of 8 bps/Hz. The other possibility is that each subband uses the same 32-AM alphabet with 32 states on a single real axis. With this set up, the actual channel data rate is $58 \times \log_2(32) \times 2 \times 100,000 = 58$ Mbps, the actual bandwidth efficiency is $58 \times 2 \times \log_2(32) / 60 = 9.7$ bps/Hz, which is about 3% less than the theoretical bandwidth efficiency of 1024 QAM, and could be more than the actual bandwidth efficiency of 1024 QAM if the 10% additional bandwidth required for pulse-shaping filter is considered.

Our observation from the example is that based on the first characteristic of the SDM, no pulse-shaping filter is needed at the transmitter; hence SDM can achieve higher actual bandwidth efficiency than equivalent QAM.

In view of the example derivation above, we have

$$\text{SDM actual channel bandwidth efficiency} = \frac{\text{Number of active subbands} * \text{subband bandwidth} * \text{bandwidth efficiency of each subband}}{\text{Channel bandwidth}} = \frac{\text{Number of active subbands} * \text{bandwidth efficiency of each subband}}{\text{Number of active subbands} + \text{Number of guard-bands}} \quad (9)$$

Using (9) and noting that the number of the guard-bands is always 2 regardless of the channel bandwidth due to the extremely high stop-band attenuation of the SDM subbands, we can see that the actual bandwidth efficiency will get better when the number of active bands increases with any increment of the digital cable channel bandwidth (e.g. 12 MHz or 18 MHz).

Thermal noise performance for SDM can be calculated and simulated by Additive White Gaussian Noise (AWGN) model. Since the theoretical AWGN performance for a SDM subband with S -AM alphabet ($S = 2, 4, 8, 16, 32, 64, \dots$) is equivalent to a S^2 QAM, so 1SDM, ..., 5SDM and 6SDM with corresponding 2-AM, ..., 32-AM and 64-AM alphabets have the same AWGN performance as QPSK, 16 QAM, ..., 1024 QAM and 4096 QAM, respectively [5].

Bit-true simulation results and lab-measured data on the baseband prototype system in Figure 2 show that both BERs are close to the theoretical values.

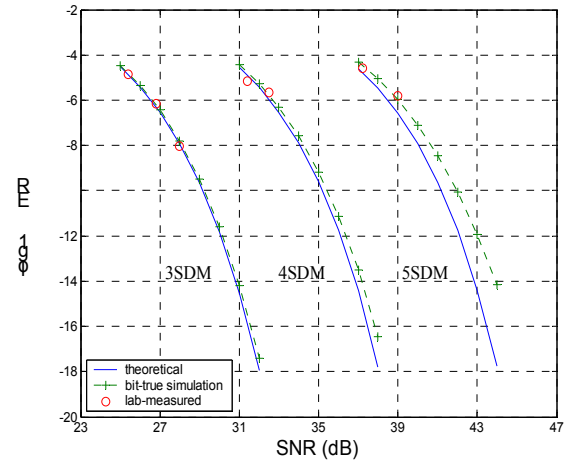


Figure 2: AWGN BER performance of SDM

Noting that the choice of alphabet for each subband is independent of other subband, we can in fact choose the alphabet for each subband differently. The freedom of doing so allows us to set the bandwidth efficiency/modulation density for each subband according to the channel condition, thus finely tune and optimize the trade off between error performance, channel signal to noise ratio and over all bandwidth efficiency. As we will see in the next section, the ability of changing bandwidth efficiency on a fine frequency scale instead of on a whole channel scale will help SDM providing higher tolerance to composite distortions.

SDM channel also has higher capacity/cost ration than combination of logical QAM channel. One of the major reasons is that SDM receiver has less complexity than comparable single QAM receiver for the fifth characteristic of SDM. In the final version of the paper, we will give more detailed comparisons between SDM and logical combination of QAM channels.

IV. SDM MITIGATES COMPOSITE DISTORTIONS

Composite distortions are produced by amplifier nonlinearity caused intermodulation of analog TV carriers. The dominant components of the distortions are Composite Triple Beats (CTB) and Composite Second Order (CSO) [7] [8]. These distortion components typically have average power levels 12~15 dB below the thermal noise level. Even though their low average power levels appear to be harmless, due to their statistical properties, the random peak envelope power can be significantly higher to cause large performance degradation for 256 or higher order QAM [8].

Other than asking operators to control and improve CSO/CTB levels through carefully choosing channel frequency offsets and maintaining head-end transmitter aggregate noise power at low levels, etc., the main methods to mitigate composite distortion at the baseband digital modulation level that have been proposed include increasing interleaver depth and improving adaptive equalizers [8]. Just by looking at the performance data, these two methods appear to be adequate. However, if we look closely, there are problems associated with each of them.

For longer interleaver depth, first we know that the prices of increased interleaving depth include increased latency, which affects the quality of service in another way. Second some longer interleaving depth required to handle CSO/CTB transients are not even supported by lots of set-tops.

For improved adaptive equalizers, according to [8], it is possible to have adaptive equalizers to converge to a state

that forms a sharp notch at the interferer frequency. However, to achieve a sharper notch through adaptive equalization, higher number of equalizer taps is required. Naturally, higher number of equalizer taps requires more demodulator complexity and more system throughput delay. In the case of non-blind equalization, which uses a training sequence to obtain the optimal equalizer taps, higher number of equalizer taps also requires longer training sequence. The longer training sequence again causes more throughput delay and overhead.

Adopting the SDM approach for baseband modulation will achieve actual bandwidth efficiency higher than 256 QAM and obtain inherent capability of mitigating composite distortion effect. In the event of excessive composite distortions, as we will discuss in the following, SDM does not need to fall back to a lower bandwidth efficiency mode completely unlike 256 QAM has to fall back to 64 QAM for the entire channel. Instead, SDM can fall back to a lower modulation density only at the subbands being affected by the composite distortions most severely, thus allowing reliable service without significantly lowering the bandwidth utilization. In addition, SDM is not contradicting with those proposed improvement done by the operators or proposed longer interleaving depth. When these improvements are available, SDM can work with these methods to provide an even higher tolerance to the composite distortions. When the limits of the HFC plants or other restrictions render these methods unusable, SDM alone still can provide a more reliable service.

To fully understand the effect of composite distortions on SDM, we first look at some properties of CSO/CTB.

Since CSO/CTB are produced by intermodulation of analog carriers, they are of narrow frequency nature. The typical power bandwidth of an individual beat is about 10- 20 kHz. Remember that the bandwidth of each subband of a SDM channel is a design choice, we can choose it to be convenient for overall system requirements yet wider than a typical composite beat component. In the example of Section III, we used 100 kHz as the subband width. It is indeed wider than the typical bandwidth of a composite beat component.

Another important property of the CSO/CTB distortions is that the locations of all the beat components can be calculated [1], [9]. Using the information of the beat components locations, we can check the composite distortion locations within a 6 MHz cable channel. Calculation shows that with 100 kHz wide subbands, only 14% of sub-bands will experience direct beats products. Due to the narrow bandwidth nature of the composite distortions and high independency of SDM subbands, only the two immediate neighboring subbands will be affected, thus the maximum number of potential neighboring subbands to be affected will be no more than 24% of the overall subbands, leaving a minimum 62% of sub-bands clear of any CSO or CTB beat impact [5]. So even without protection of the error correction codes and/or interleaver, in the event of excessive composite distortions, we can lower the bandwidth efficiency on those 38% affected subbands without changing the 62% unaffected subbands. As an example, we can lower the bandwidth efficiency of the 38% affected subbands from 256 QAM equivalent 8 bps/Hz to 64 QAM equivalent 6 bps/Hz to maintain the reliability, while leaving the rest 62% subbands still at 8 bps/Hz. The end results is that the overall channel capacity is about

$6/8 \times 38\% + 8/8 \times 62\% = 90\%$ of the normal channel capacity, but still gain better reliability.

To look further at the error correction and interleaving protection, we note that the duration of a random composite distortions pile up burst is often inversely proportional to the distortion power bandwidth [8]. Given a 10 kHz power bandwidth, a distortion burst can be 100 μ s long. Since these bursts are highly localized in subbands, they only affect the symbols in one subband. Recall from Section II, the interval between two consecutive symbols in the same subband is 5 μ s for 100kHz wide subbands. So a distortion burst of 100 μ s only covers about 20 symbols in the same subbands. The fact that a distortion burst only affects a low number of symbols implies that an appropriately chosen Reed-Solomon (RS) type error correction coding can correct most errors caused by composite distortions bursts. In the final version of the paper, we will have more detail about the effect of RS code on SDM and we will also show that a simple standard length interleaving will have the similar results as well.

V. TEST RESULTS OF A SDM PROTOTYPE SYSTEM

A Broadband Physics, Inc. SDM prototype system has been tested over RF channel with some simulated composite distortions. The system modulation density was set at 3SDM, which according to the previous sections provides a theoretical 64 QAM equivalent bandwidth efficiency of 6 bps/Hz. The actual system bandwidth efficiency is 5.8 bps/Hz (see Section III). The 6 MHz channel occupies RF spectrum of 582~588 MHz and is centered at 585 MHz. An approximate “spread” interference tone was generated with an occupied bandwidth of 20 kHz. This interference tone

was modulated at different frequencies within the channel. The overall uncoded channel BERs were measured under different levels of interference tone and Additive White Gaussian Noise (AWGN) at constant level of -37dBc . Parts of the results are shown here in the chart below.

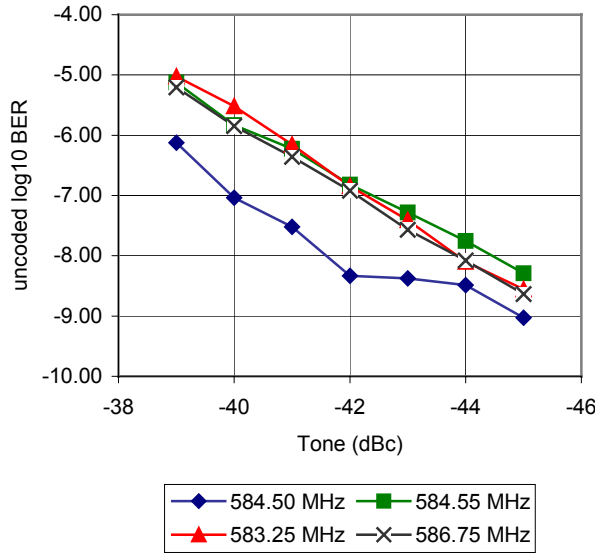


Figure 3: BER Performance of 3SDM over RF Channel with AWGN and Narrowband Interference

The most interesting feature of the chart is that for interferences with the same power level, their effects are also dependent on their locations within the channel. This distinctive feature, which is not available with single band modulations, is a direct result of the multi-band approach and the subband independency of SDM.

Recalling that the 6 MHz RF channel is divided by 60 SDM subbands, each subband occupies 100 kHz bandwidth, we find that three of the four interferences, namely the tones at 583.25, 584.55 and 586.75 MHz, are in the middle of a subband and the other one (at 584.5 MHz) is on the boundary of two neighboring subbands. The chart shows that the interference on the boundary of two

neighboring subbands causes less BER degradation than interferences falling in the middle of a subband. Because of the SDM subband independency, the increased bit errors due to the narrow band interferences are from one and two subbands, respectively, for the interferences falling in the middle of a subband and on the boundary of two neighboring subbands. Let us assume that the average number of increased errors caused by an interference with a certain power level falling in the middle of a subband is Ed . Another interference with the same power level but falling on the boundary of two neighboring subbands will actually have half the interfering power in each affected subband, or equivalently, 3 dB higher signal to interference ratio. In view of this, the average number of increased errors in both affected subbands will be $Ed/2 \ll Ed$, and the overall number of increased errors will be $2*Ed/2 \ll 2*Ed/2 = Ed$. Hence, the overall BER degradations caused by narrow interferences falling on the boundary of two neighboring subbands are less than degradations caused by those falling in the middle of a subband.

The simple analysis above is not valid for the single band modulation or multi-band modulation without subband independency. In those cases, error degradations are not determined by the interference power within an individual subband or a spectrum subsection, so the type of results in Figure 3 can only come from multi-band modulation schemes, with highly independent subbands, such as SDM.

In the near future, we will present more test results to further verify the CSO/CTB mitigating capability of SDM.

VI. CONCLUSION

In this paper, we have introduced the basic concept of Sub-band Division Multiplexing (SDM). Several essential characteristics of the SDM have been presented. As a digital modulation scheme, advantages of the SDM associated with its characteristics were also discussed. These advantages include:

1. Effective bandwidth efficiency is improved, as no pulse-shaping filter is needed.
2. Transmitting can be tailored easily to fit under transmission masks since sub-bands are highly independent, nearly orthogonal, and can be turned active or inactive easily.
3. Transmitted signal can be made real to have higher phase noise resiliency.
4. Less receiver complexity can be achieved, as no time domain equalizer is needed.

In the latter part of the paper, we focused on the application of SDM in cable channels. We discussed that due to its characteristics, SDM can provide higher actual bandwidth efficiency than QAM with the same theoretical bandwidth efficiency. We further discussed that also due to various advantages of SDM, especially its subband independency, SDM can effectively mitigating composite distortions caused by CSO/CTB, thus relax the constraints on the RF channels and lessen the burdens on operators.

In Section V, we also presented the test results for a prototype SDM modem to verify the theoretical results.

ACKNOWLEDGEMENT

The authors wish to thank Mr. Steve Anderson and Mr. Rama Nagurla for their effort in constructing and testing the Broadband Physics, Inc. SDM prototype system and generating the test results.

REFERENCES

1. M. E. Laubach, "Moving towards Shannon's Limit, Sub-Band Division Multiplexing could help MSOs pump up existing cable", CED Magazine, Reed Business Information, September 2003.
2. W. J. Miller, U.S. Patent # 5,367,516, "Method and apparatus for signal transmission and Reception", November 1994.
3. G. Cherubini, E. Eleftheriou, S. Olcer and J.M. Cioffi, "Filter bank modulation techniques for very high speed digital subscriber lines", IEEE Communications Magazine, May 2000.
4. G. Strang and T. Nguyen, "Wavelets and filter banks", Wellesley-Cambridge Press, 1997.
5. Broadband Physics, Inc, "SDM Technology", www.broadbandphysics.com, 2001-2004.
6. E. A. Lee and D. G. Messerschmitt, "Digital Communication", 2nd Ed. Kluwer Academic Publishers, 1994.
7. W. Ciciora, J. Farmer, D. Large and M. Adams, "Modern Cable Television Technology", 2nd ed., Morgan Kaufmann, 2004.
8. R. D. Katznelson, Statistical Properties of Composite Distortions in HFC Systems and Their Effects on Digital Channels, Cable-Tec Expo Proceedings, SCTE, 2002.
9. T. B. Warren and J. Kouzoujian, "Some notes on composite second and third order intermodulation distortions", Matrix Technical Notes, MTN-108, Matrix Test Equipment Inc., 2003.

THE INTELLIGENT NETWORK: DYNAMICALLY MANAGING BANDWIDTH AT THE OPTICAL LEVEL

Gaylord A. Hart and Steven Robinson
Mahi Networks, Inc.

Abstract

CATV network bandwidth requirements are evolving at a rapid pace, driven by deployment of new services, increasing penetration of existing services, and the ongoing transition from analog to digital services. As services migrate to IP based and on-demand content delivery, bandwidth requirements vary dynamically in real time, and a minimum QoS is required to ensure adequate service performance. All of these place stringent requirements on the network.

Managing evolving & dynamic bandwidth requirements is complex, but emerging reconfigurable optical add/drop multiplexers (ROADMs) allow transport bandwidth to be managed effectively at the optical layer. Moreover, ROADMs enable fully automated optical transport systems which eliminate design and cut-over errors, accelerate service delivery, and lower network costs. This paper explores evolving services, optical transport technologies, and bandwidth management mechanisms in the context of migrating to an all-digital, IP based network that lowers both CapEx and OpEx costs.

INTRODUCTION

Bandwidth requirements in CATV networks are rapidly changing, driven by deployment of new services (VoIP and HDTV), increasing penetration of existing services (DTV and cable modems), and the ongoing transition from analog to digital and from circuits to packets. This picture is further complicated by the transition to on-demand services (with large variations

between peak and average usage) and the migration to IP based services (which are connectionless, and therefore sometimes unpredictable in traffic requirements). Layer onto this the need to maintain QoS across a wide range of services, and bandwidth management becomes even more critical to tomorrow's cable network.

There are two key components to managing bandwidth in the future: the naturally changing bandwidth requirements as the network and service penetration evolve over a period of time and the real-time bandwidth management requirements imposed by everything-on-demand (EOD) and the converged IP network. To minimize long-term CapEx and OpEx costs (and hence remain competitive), MSOs must plan for these future requirements today and deploy systems that evolve with the network without forklift upgrades and without compromising service quality.

The basic HFC architecture is a long way from running out of bandwidth: increasing bandwidth in the future is relatively easy and inexpensive to accomplish simply by subdividing optical nodes (i.e., reducing the number of homes served per node and thereby increasing the bandwidth available per home), and the equipment for doing so exists today and is relatively inexpensive. Alternatively, a significant amount of network bandwidth can be made available by converting the broadcast analog TV services on the network to digital services (roughly a ten to one savings in bandwidth).

Because the HFC access architecture largely relies on transparent transport pipes,

deploying new services on the network frequently does not require any upgrades to the HFC network itself, thus greatly lowering new service deployment costs and shortening time to market. These new services do, however, place new requirements on the transport and switching components in front of the HFC plant.

Over the next few years, the greatest changes in the CATV network will occur in front of the HFC access plant at those points in the network where services are aggregated, switched, and transported in the purely digital domain. For the MSO, this includes ISP and telephony POPs, content storage and origination points, regional and metro headends, primary and secondary hubs, and large businesses where services will be delivered directly via fiber.

A network's architecture is bounded by the optical transport paths connecting it, both physically and from a bandwidth perspective. These paths can be a bottleneck to delivering sufficient bandwidth and QoS as the network evolves. While it is common to engineer these paths to provide sufficient bandwidth under a pre-defined set of conditions, this can be costly, either requiring constant re-engineering to add incremental bandwidth as needed or resulting in stranded bandwidth if current traffic loads are far less than the deployed transport bandwidth.

This problem only worsens as services migrate to content on demand and IP based delivery. Both of these natural evolutionary steps, though very efficient in only requiring bandwidth when it is actually needed, result in bursty traffic, which makes traffic engineering even more complex. Needless to say, network complexity and traffic variability are only increasing, and new approaches are required for network design and operation in the future if we are to ensure

the lowest network costs while maintaining an acceptable QoS.

Network engineers are used to designing optical transport paths as static circuits providing dedicated bandwidth to highly predictable traffic. As new CATV services are deployed, as penetration increases for existing services, and as services migrate to connectionless delivery via IP, new demands are being placed on the transport network and transport engineering. At the same time, rising customer expectations for service reliability are requiring more redundancy in the network. To keep up with these demands and to address emerging requirements for real-time bandwidth management (brought on by EOD and IP services), it is necessary the detailed optical transport network design, configuration, migration, and operational processes be automated, essentially masking much of the network complexity from the engineering and operational processes while enabling accelerated network evolution.

In its fullest implementation, such an optically reconfigurable network would dynamically modify itself in real time to respond to changing network conditions and service requirements. This would be analogous to the way a router automatically discovers paths through the network and dynamically routes packets to their destination based upon changing network conditions and traffic requirements. It now begins to make more sense to treat the optical transport layer (Layer 1 in the OSI model) as an extension of the switched layers residing above it, supporting integration of bandwidth management and QoS across the optical, Ethernet, and IP layers.

All of these objectives for optical transport automation and a dynamic optical layer can be accomplished with ROADMs using DWDM and GMPLS. When these

technologies are coupled, an underlying optically reconfigurable network is possible which flexibly and dynamically supports future network evolution as well as bandwidth and content on demand, and without having to manually re-engineer and re-configure each component of the network when changes are made.

Presently, dense wave division optical add/drop multiplexers (DWDM OADMs) allow MSOs to collapse multiple parallel service transport networks onto a common optical network that supports multiple protocols and services over a single fiber. With modern network planning tools, the design of both the optical and service layers may be fully automated. This simplifies and speeds up network design, eliminates the need to memorize and understand complex design rules, and reduces design errors.

OADMs, however, typically rely on fixed configurations and components which are determined at the time of the initial design. This restricts the degree of automation which may be applied in the future to network evolution and slows down the upgrade process itself. ROADMs, because they allow their configuration to be controlled remotely via software, enable provisioning and operations to also be fully automated. This includes network topology discovery and service turn-up, as well as the dynamic monitoring and setting of network parameters for optimized operation. Modern ROADMs already support these capabilities today.

But as optical switching speeds increase and the GMPLS control plane becomes more content-aware and tightly coupled to Layers 2 and 3 above it, ROADM based CATV optical transport networks will become even more powerful, being able to dynamically re-configure transport paths and bandwidth in

real time based upon changing service and content requirements, as well as subscriber on-demand service usage.

ROADM TECHNOLOGY

OADMs have been deployed widely in optical transport networks over the last few years. Most of these rely on fixed wavelength components (lasers, multiplexers, filters, etc.) which require a significant amount of manual network design as well as manual configuration and provisioning. A new class of fully reconfigurable OADMs has recently emerged which enables automation of these processes as well as automated and dynamic network operation. These ROADMs are built around flexible optical components which can be controlled via software and an intelligent control plane which supports process automation. The primary underlying technologies which define ROADMs are outlined below.

ROADM Architectures

Several architectural alternatives, based on a variety of markedly different optical technologies, exist today. Early ROADM technology, based on discrete optical-mechanical switches, filters and variable optical attenuators (VOAs), is shown in Figure 1, below.

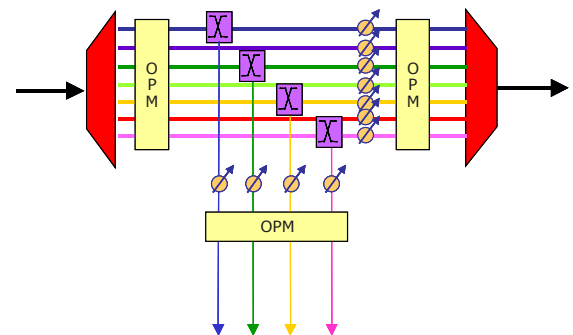


Figure 1. Discrete ROADM Architecture

While simple to implement, since most of the technology is commercially available, this approach utilizes many discrete optical components. The result is very high insertion loss, very high cost, and large size—thus preventing its widespread acceptance.

The second architecture to be used for ROADMs, the wavelength blocker architecture, is shown in Figure 2, below. Essentially, this design splits the incoming DWDM signal into a drop and through path. An integrated DWDM demultiplexer (Demux), VOA and multiplexer (Mux) form the core of the wavelength blocker. Typically blockers are implemented using Micro-Electrical Mechanical Systems (MEMS) or Liquid Crystal Display (LCD) technologies.

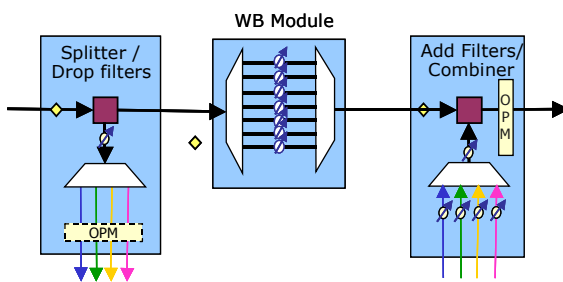


Figure 2. Wavelength Blocker Architecture

While this architecture reduces the number of discrete components, it forces the MSO to pay for all wavelengths at each node on day one. Furthermore, this design only manages the through wavelengths—not the add or drop wavelengths. In fact, most implementations of this architecture use inflexible fixed filters for the add and drop wavelengths. Hence this design is actually a “semi-reconfigurable” ROADM.

A variant of the wavelength blocker architecture is the combined blocker / adder architecture shown in Figure 3, below. This

design integrates the add Mux with the wavelength blocker. This design eliminates the extra add Mux, but at the expense of requiring additional optical switches.

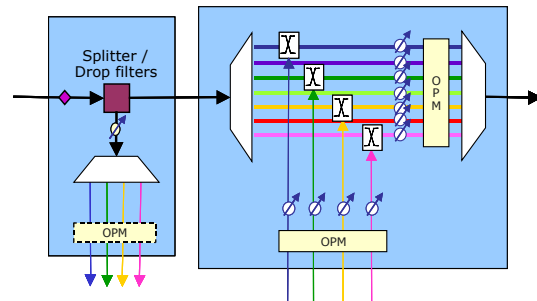


Figure 3. Combined Blocker / Adder ROADM

To be cost effective, this design typically integrates all of the filters and switches into a single module (similar to the wavelength blocker architecture). However, this again forces the MSO to pay for all wavelengths on day one. While this design now manages add wavelengths, it still does not manage drop wavelengths. Furthermore, this design now permanently locks the add wavelengths to a fixed wavelength design—thereby making fully tunable lasers only good for sparing (no dynamic wavelength management). Thus, this design is also a “semi-reconfigurable” ROADM.

Planar Lightwave Circuits (PLCs) are one technology used to implement this combined blocker / adder design. While they have the potential to integrate complex optical components onto one or more substrates (such as silicon), manufacturing yield and power management are still a challenge.

More recently, a new architecture based on Multi-Port Wavelength Selectable Switches (MP-WSS) allows for completely reconfigurable ROADM functionality. This architecture is shown in Figure 4, below. Wavelength switches have the ability to

direct one or more wavelengths from an incoming DWDM signal to one or more output ports (usually with individual VOA-like power control for each wavelength).

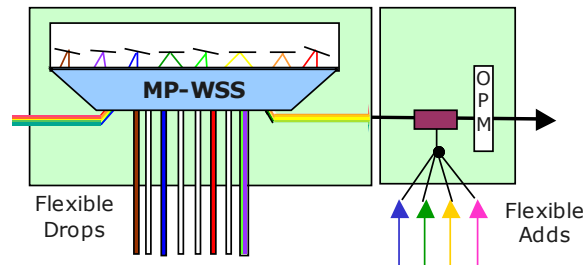


Figure 4. Fully Flexible WSS Based ROADM

Note that in this architecture, the MP-WSS is located on the drop side. In this configuration, the ROADM can manage both the drop and through wavelengths. It can also direct multiple wavelengths to a single drop port for low cost optical ring interconnection or physical mesh nodes. This design is inherently more efficient because the same demultiplexer is used for both through and drop wavelengths.

Another interesting aspect of this design is the use of broadband (wavelength independent) optical add ports. In this design, dynamic wavelength provisioning with full C-band tunable lasers is now possible. MP-WSS technology enables a “fully reconfigurable” ROADM architecture.

Optical Backplane

With the advent of the fully flexible ROADM (using the WSS), practical optical backplanes can now be implemented on an OADM chassis. Because wavelength selective switching eliminates the need for manually configured optical jumpers, an optical backplane permits the 100 or more

manually placed front panel optical interconnect jumpers to be moved inside to the rear of the chassis. Reciprocal optical connectors on the back of each transponder allow these optical interconnects to be made when the transponder is plugged into the chassis. Without the WSS, however, an optical backplane would require all slots to use pre-defined wavelengths and thereby waste precious rack space.

An optical backplane frees MSOs from the tangled mess of jumpers interconnecting the lasers & receivers to the DWDM filters and amplifiers and from the need to carefully map and record these interconnects. Operators simply connect the client-side service fibers and walk away.

Tunable Lasers

The recent advent of widely tunable lasers has enabled an equivalent capability for full reconfigurability on the transmit side of the ROADM. Initially, MSOs were interested in tunables primarily for sparing purposes. This typically yields a 32:1 savings in spares and provides tremendous cost savings. More recently, MSOs have realized that tunables also substantially reduce delivery lead times and allow for significant equipment reuse.

More significantly, however, widely tunable lasers, when coupled with wavelength selective switches with single lambda add/drop granularity, enable greatly simplified provisioning as well as dynamic wavelength management. Combined with an optical backplane, remote provisioning and management are also possible. Operators no longer need to match fixed wavelength lasers with fixed wavelength filters using manually configured jumpers during installation. Instead, operators can simply point and click to remotely provision or re-provision

wavelengths. This capability delivers tremendous OpEx and CapEx savings.

Optical Power Monitoring and Control

Initial OADM implementations provided only broadband, non-real-time optical power monitoring of the aggregate DWDM signal. However, to properly set up, manage, equalize, and optimize wavelengths, real-time optical power monitoring (OPM) and power level control are required for each lambda at each node. Real-time, direct measurements are more reliable and provide for rapid fault identification as well as support for newer protection schemes such as shared optical protection switching. Ideally, and for a fully robust system, all optical inputs, outputs, adds, and drops would be monitored in real time at each node for system optimization and full fault isolation.

Robust Variable Gain Amplifiers

Early DWDM amplifiers required operators to operate them in fixed gain mode. To ensure a flat gain profile, operators had to manually “pad-out” every span. This unnecessarily added noise to the DWDM signal—limiting reach. As new wavelengths were added or span losses changed, the system ran the risk of becoming unbalanced. Furthermore, any upstream wavelength changes (such as fiber cuts) could cause downstream errors (on both working and protection wavelengths).

Newer DWDM systems are now being implemented with variable gain amplifiers with transient control. Variable gain amplifiers can provide continuous, automated broadband gain adjustments to correct for slow changes in span losses or a change in the number of DWDM channels.

For rapid changes (such as a fiber cut), ultra-fast amplifiers can now provide robust transient control by quickly clamping the gain change to less than 0.3 dB in less than 1us. This accuracy is needed since the transient is additive with cascaded amplifiers. Combined with the ROADM capabilities already outlined above, both per-channel & broadband optical power level management can be fully automated and dynamically controlled in real time, thereby simplifying optical layer design and network management and operation.

GMPLS Control Plane

One final technology is required to fully automate the provisioning—an integrated control plane that supplies the intelligence and communication needed to control the ROADMs in the network. More recently, DWDM vendors have begun implementing GMPLS (Generalized Multi-Protocol Label Switching) control planes. Typically a dedicated wavelength carrying an Optical Supervisor Channel (OSC) is used to communicate the GMPLS messages between nodes. This wavelength is completely independent of the service bearing wavelengths and also supports provisioning communications from the Craft Interface and Element Management System (EMS).

GMPLS enables DWDM nodes to seamlessly work together in a network to provide resource and inventory auto-discovery, topology information, service setup, signaling, path computation (including Routing and Wavelength Assignments), light path setup, and management. Following industry standards, a GMPLS control plane enables true cross-network, A to Z provisioning.

BANDWIDTH MANAGEMENT

When wavelength selective switching, tunable lasers, and an optical backplane are combined in a ROADM along with a GMPLS control plane, a powerful platform is created with an extensive set of capabilities. Because of these capabilities, ROADMs are much more flexible in supporting service migration and network evolution and in automating these. For the following discussion, this paper assumes a ROADM has all the attributes listed above (i.e., genuinely is reconfigurable), though this is not always the case.

Network Design

Bandwidth management in the optical transport network begins with the initial network design. Limitations built-in to the initial design can limit future growth and network evolution, or at least make them much more painful and costly to achieve. Ideally, MSOs should plan for one-time network engineering, which allows unrestricted and non-disruptive addition or removal of services, wavelengths, and nodes to the network at any time.

The initial system design, for example, should readily support the migration (or mix and match capability) from 2.5G to 10G to 40G wavelengths without re-engineering and upgrading the network at each stage. Similarly, the ability to add or delete nodes in the network without re-engineering amplification or dispersion compensation is a significant advantage. ROADMs, because they are reconfigurable and can dynamically adapt to new network configurations, make this possible. ROADMs also further simplify and lower the cost of network evolution because they are remotely configurable, which in most cases eliminates the need to

roll a truck to each node when an upgrade is performed.

Many modern software-based network planning tools allow the optical layer design process largely to be automated. This is a critical component in enabling flexible bandwidth management in the future. A good planning tool also allows the MSO to flexibly control design options (if necessary) and to optimize the design based upon selectable criteria such as cost or wavelength conservation.

However, if an OADM uses fixed filters or other components which are not dynamically compensated for by the network itself, these must manually be entered into the network design tool during the initial design and may require subsequent re-engineering of the network when upgrades are performed. By eliminating many of these fixed components or by being able to dynamically compensate for them, ROADMs simplify the initial network design process and any subsequent upgrades.

Network planning tools should also support automated design of the service layer as well, allowing an MSO to select protection options for common equipment and individual wavelengths, service and transponder types, client interfaces, and any other number of options. Once these are selected, the planning tool will automatically design the optical and service layers. At the optical layer, OSNR and dispersion compensation budgets are created, and the planning tool selects appropriate optical amplifiers and dispersion compensation modules to ensure one-time network engineering. At the service layer, common equipment and transponders are selected and assigned to specific chassis slots.

A well-designed planning tool, once it has provided a network design, will also provide detailed network drawings, span information, chassis drawings with slot assignments for each shelf component required, a bill of materials with part numbers, and even costs. The planning tool should also support exporting these items to an external spreadsheet for subsequent analysis, manipulation, or record keeping. Because the planning tool can reduce the time required to design a network to a matter of minutes, it is easy to create different network scenarios for comparative evaluation, allowing the MSO to optimize the network design based upon other service or architectural factors.

ROADMs, providing their underlying components have been designed to provide fast enough switching times, also enable more flexible protection options. Because wavelength selective switches and tunable lasers allow receive and transmit wavelengths to be assigned on the fly, protection wavelengths need not be assigned until an actual failure occurs. Such protection allows more than one path to share a common protection wavelength, which means more protected services may be placed on a given fiber. This conserves fiber while still providing complete service protection.

Good bandwidth management also includes conservation of network resources and performance optimization of these resources. In other words, you should use as few resources as necessary to get the job done, and you should get the most out of them. A well-designed ROADM conserves fiber and wavelengths on a fiber by eliminating stranded bandwidth caused by wavelength banding and by enabling unrestricted wavelength assignment and reallocation (single lambda granularity).

ROADMs which can dynamically optimize performance and optical parameters at each node also allow the typical engineering design rules to be relaxed, which enables transport over longer distances and through more nodes. All of these provide the MSO better bandwidth utilization and at a lower cost.

Collapsed Service Transport Networks

Most CATV networks and the services delivered over them have evolved so rapidly over the last few years that many MSOs have deployed parallel networks optimized for each service. While this made economic and technical sense at the time, it is not a sustainable model for a competitive market because it is costly to maintain, manage, operate, and upgrade these networks and because it makes inefficient use of network resources.

For most MSOs, the first stage in network evolution is collapsing these parallel networks onto common network elements and infrastructure wherever possible. Multi-service, multi-protocol ROADMs are ideal for collapsing the transport sections of these parallel networks onto a single, unified DWDM network (typically a ring) which provides fiber relief and service protection. Because ROADMs enable non-disruptive upgrades, follow-on network evolution can be transparently accomplished as new services are added and old services removed. Because ROADMs also enable flexible add/drop wavelength assignment, lambdas may constantly be recycled with service changes, allowing full utilization of existing fiber resources.

ROADMs with a GMPLS control plane offer another significant benefit, as well. Because GMPLS based ROADMs combine reconfiguration capabilities, intelligence, and

communication between nodes, they also support automated network topology and inventory discovery. This eliminates the need to manually provision these into each network node and the EMS. At the same time, this capability allows the network to automatically turn up the individual wavelengths and services to be carried over it.

Another key aspect of bandwidth management for these unified networks is the ability to deploy only that bandwidth which is needed to support today's services, but at the same time to provide migration capacity and capability for tomorrow's bandwidth requirements. MSOs should not only expect a ROADMs to support mixing and matching 2.5G and 10G wavelengths, but 40G wavelengths as well. While demand for 40G wavelengths may not be strong today, it will be in the future. In keeping with the one-time network engineering philosophy outlined above, this should be accomplished transparently, that is, without any additional upgrade costs or new restrictions on span budgets, node counts, or ring circumference. This pay-as-you-grow approach optimizes CapEx by tying network costs to service revenue, but it also eliminates or puts off into the future costly forklift upgrades and the deployment of new fiber.

Transport Network Evolution

Once any parallel single-service transport networks have been collapsed onto a common multi-service DWDM optical network, the next stage of network evolution becomes a matter of supporting bandwidth growth, service conversion, and new service deployment, as well as ensuring sufficient QoS is provided for all these services.

Bandwidth growth and new service deployment are relatively easy to provide

with an OADM simply by adding transponders, providing sufficient lambdas are available to support the new transponders, or by migrating from 2.5G to 10G to 40G transponders, providing the network and the OADM have been designed to support these. ROADMs configured for one-time network engineering are ideal for supporting this type of growth. New transponders may be installed without any additional network engineering, and the ROADMs will dynamically adjust to the new traffic load. And unlike many OADMs with fixed serial filters or multiplexers, the ROADMs permits transponders to be added without any service disruptions.

Seamless service conversion or migration is also readily supported by ROADMs. A good example of this is the migration from TDM voice to VoIP. In this case, new wavelengths are deployed to support the emerging IP service, while existing wavelengths continue to be used to support the legacy TDM service. As customers migrate to the new service, the legacy DWDM wavelengths may be torn down and recycled on the transport ring for new services. ROADMs enable flexible and unrestricted wavelength recycling and non-disruptive migration for existing services and customers.

As part of the bandwidth management role, OADMs may also be used to help manage QoS. Some MSOs have elected to segregate at various points in their networks, at least initially, services which have different QoS requirements and traffic characteristics. For example, IP data services and VoIP services, though both are IP services, have very different QoS requirements with respect to latency, jitter, and lost packets and very different traffic characteristics with respect to average to peak bit rates. It is possible to segregate

these services in the network to maintain better control over QoS for each.

OADM's can be used effectively here for parallel transport of each service on a separate lambda over a single fiber, allowing unique bandwidth and protection options to be applied to each service. OADM's also typically support multiplexing transponders, which carry many multiplexed tributaries over a single wavelength and can therefore support this type of service segregation, as well. Of course, in the long run, the efficiencies of common transport and switching of all these services will require convergence. ROADMs can be of particular value here by enabling flexible and non-disruptive changes in the network to support dynamic service requirements and also the future transport convergence of these separate services when suitable QoS mechanisms (such as MPLS) are deployed in the network to guarantee QoS under worst-case conditions.

Converged IP Transport Network

A significant amount of equipment has already been deployed in CATV networks to support existing services and will not go away in the near future. Nevertheless, voice, video, and data services are destined for IP delivery. Indeed, data services in CATV networks are already there, with VoIP rapidly emerging, and IP video to follow.

As this transition occurs, MSOs will continue to operate parallel legacy service networks, as already outlined above, until such time comes when all customers are transitioned to IP based services and the legacy equipment is taken out of service. Multi-service, multi-protocol ROADMs can provide significant help in easing this transition while still allowing all services to be delivered over a common transport

network. As IP services begin to dominate, legacy wavelengths can simply be recycled as no longer needed, freeing up additional transport capacity for IP services. However, significant challenges remain in deploying a fully unified IP network.

IP is a Layer 3 protocol and still requires transport by a Layer 2 protocol riding over a Layer 1 Physical Layer in the network. While most MSOs have deployed some ATM and SONET in their networks, it has typically not often been a significant amount, and MSOs are now focused on Ethernet for Layer 2 and Optical Ethernet for Layer 1. OADM's can readily and simultaneously transport all these protocols and services and can ease the transition as MSOs migrate their networks.

Ethernet is a wise choice for building CATV networks. Ethernet now scales effectively and inexpensively across LAN and WAN environments, is very easy to deploy and manage, and as a connectionless service is well-suited for IP, which is also connectionless. Recent Ethernet standards for virtual LANs (VLANs) make Ethernet even more powerful and support QoS and traffic prioritization capabilities in Ethernet for the first time.

When coupled with multi-protocol label switching (MPLS), which allows even tighter coupling between the Ethernet and IP layers, Ethernet clearly offers advantages that cannot be found with ATM or other protocols. Similarly, optical Ethernet, when coupled with GMPLS, offers advantages that cannot be found in SONET, which though widely deployed and quite capable of transporting Ethernet, is still a TDM transport technology.

As legacy services fade away and Ethernet and IP dominate, more wavelengths on the optical transport network will be dedicated to

Ethernet. While services may initially be segregated for transport for QoS reasons, in the long run this will not scale effectively, especially in light of modern VLAN Ethernet switches which also support MPLS and can prioritize traffic. So the network will not only migrate to Ethernet transport, but to multi-service transport in each pipe with each service provided sufficient transport bandwidth to ensure reliable performance.

Many benefits are still to be derived from the over-subscription that statistical multiplexing allows, especially when larger pipes can be used for transport of a larger number of services. So one can also expect a migration to wavelengths with ever more transport capacity (2.5G to 10G to 40G per lambda). Ultimately, it is also more cost effective to use fewer (but higher capacity) switch and router interfaces. MSOs should therefore, as they make transport decisions today, look for solutions that provide one-time network engineering and transport capability for 40G wavelengths.

As CATV networks migrate to IP based services, a parallel migration to on-demand services is also taking place. Broadband data access and VoIP services are already demand based services, for the most part only consuming bandwidth when the services are actively being used. Video on demand (VOD) is also widely deployed, but most CATV subscribers still receive their video services via broadcast analog or digital TV.

On-demand services have the benefit of only consuming network bandwidth when a service is requested by a subscriber, thus allowing statistical multiplexing at the service level. However, on-demand services can also result in a very high peak to average bandwidth usage ratio, which requires careful network planning to prevent network congestion and poor service delivery. This

can be particularly true for video, which not only has relatively large bandwidth requirements, but which is also sensitive to latency and lost packets.

Because on-demand usage is also tied to customer preferences and other uncontrolled variables, the peak to average ratio may also vary significantly over time. For example, television viewing typically peaks in the evening, but there can also be significant traffic peaks created by special or unpredictable events. And this can be true across voice, video, and data services. While multicasting in an on-demand environment can mitigate the effects of these peaks for video, it cannot eliminate them.

Traffic planners are faced with a difficult choice: provide sufficient bandwidth for peak demand and let a significant amount of the bandwidth go unused most of the time, or provide less bandwidth and let the network congest under peak traffic loads. The first option is expensive, and the second results in poor service. Of course, most traffic planners try to hit a reasonable compromise, but with service usage varying dynamically over time (both in real-time and as a result of demographic trends), this is not always possible. Clearly, it would be desirable to have the network intelligently monitor its own loading and dynamically apply network resources where and when needed. Such an intelligent network would allow sharing a smaller number of network resources across a wider range of applications and conditions.

ROADMs, given their flexibility and on-the-fly reconfigurability, are entirely capable of supporting such a dynamically configured network in the transport domain. Under such a scenario, network planners could deploy a sufficient amount of transport bandwidth to cover average bandwidth requirements and provide additional uncommitted wavelengths

which could be brought to bear when and where needed to prevent congestion as traffic demand builds. Load-sharing switches or routers would provide complementary bandwidth delivery functions at the Layer 2 and/or 3 levels.

Since network bottlenecks can occur in both the transport layer (insufficient bandwidth to carry the traffic) and switching layer (insufficient capacity to switch traffic), a dynamic transport network could provide additional wavelengths when transport capacity is exhausted or route existing wavelengths around congested switches when switching capacity is exhausted.

In this intelligent network, coordination between Layers 1, 2, and 3 would be implemented and automated by a GMPLS and MPLS control plane, which allows label-switched paths to be created and torn down as needed. Because the overall network status would then be known at all layers, optimization could be applied intelligently where it makes the most sense.

At the ROADM level, as congestion begins to build in the network, the GMPLS control plane would assign and route wavelengths between nodes on the ring, providing additional bandwidth in real time where and when it is needed. This mechanism could also be used to route traffic around a congested network segment or element, relieving traffic from congested or near-congested areas and applying it where sufficient resource are available to handle it. This approach can also be used to provide protection switching within the network by dynamically reassigning wavelengths and or transponders to compensate for equipment failures or fiber cuts.

Just as routing protocols dynamically allow routers to discover the most

appropriate paths for moving datagrams through a network (and to compensate for failed routers), GMPLS and its associated routing protocols can route wavelengths through the optical transport network as needed to optimize network resources and capabilities. While current GMPLS and MPLS standards do not support this degree of integration across layers 1, 2, and 3, this capability will no doubt be implemented at some point.

An intelligent network built upon these principles would require less equipment since resources are optimally applied only when and where needed and would simplify operations since network reconfiguration is done automatically and dynamically, and not manually by traffic planners. Overall, this would result in greater CapEx and OpEx savings for the MSO while providing greatly increased service reliability for subscribers.

SUMMARY AND CONCLUSIONS

CATV optical transport bandwidth requirements are constantly changing as a result of increasing service penetration, deployment of new services, and the transition to digital and IP based services. At the same time, on-demand services are rapidly being deployed and will likely displace broadcast services at some time. On-demand services significantly add to the complexity of traffic planning and bandwidth management. New demands are also being placed on the network to ensure adequate QoS is provided for each digital service to ensure reliable delivery. Service and network evolution is occurring at an ever-increasing rate, and new real-time constraints are being placed on the network which require dynamic optimization.

MSOs also face an increasingly competitive service environment which

requires reliable and low-cost service delivery, rapid roll-out of new services, and the lowest possible CapEx and OpEx costs. To achieve these objectives, MSOs must deploy equipment today that optimizes performance and efficiency on the existing network infrastructure and which offers low first-in costs, yet can evolve with the network without forklift upgrades. This equipment must also provide a high degree of automation for network design, service turn-up, and operation, which in turn will result in lower costs, more reliable services, and accelerated service deployment and network migration.

ROADMs deliver the best approach to managing changing bandwidth requirements and evolution in the optical transport network. ROADMs support automated optical and service layer design, as well as one-time network engineering. They provide automated topology self-discovery and service turn-up, unrestricted wavelength usage and reconfigurability, and automated

and dynamic monitoring and adjustment of network operational parameters for optimized performance. ROADMs are ideally suited for MSOs to meet their existing and future network transport requirements and also offer the most cost-effective solutions for optical transport, significantly lowering overall lifecycle CapEx and OpEx costs.

REFERENCES

"The Dynamic Network: Managing Bandwidth and Content on Demand at the Optical Level," Gaylord Hart and Zouheir Mansourati, Proceedings Manual: Collected Technical Papers of the 2004 SCTE Conference on Emerging Technologies.

Author Contact Information

Gaylord A. Hart
ghart@mahinetworks.com
303.910.7743

Steven Robinson
srobinson@mahinetworks.com
732.465.1000 ext. 1248

THE MODULAR CMTS ARCHITECTURE

John T. Chapman
Cisco Systems, Inc.

Abstract

The architecture working group of Next-Generation Network Architectures (NGNA) posed the question to the industry if a cable modem termination system (CMTS) architecture could be developed that could leverage edge quadrature amplitude modulation (QAM) devices that have been developed for the video-on-demand (VOD) environment. The author of this paper started working on this problem with a team of engineers in December 2003.

The resulting design was submitted to CableLabs in May 2004 and formally adopted by CableLabs in January 2005 as the baseline design for the Modular CMTS (M-CMTS) specification. The author of this paper is now serving as lead author for the Downstream External PHY Interface (DEPI) working group.

INTRODUCTION

A new architectural concept for building Data Over Cable Service Interface Specification (DOCSIS) CMTSs is underway at CableLabs. It is known as the Modular CMTS (M-CMTS). The concept is to extract the Physical Layer (PHY) out of the CMTS and locate it in a separate network element. That separate PHY network element would be an evolution of the edge QAM network element that has already been developed and deployed for the VOD market.

The motivation for this is several reasons. First, there is the promise of cheaper downstreams. The cost of a downstream on a traditional CMTS is ten to twenty times the cost of a downstream on an edge QAM device. This is mainly due to the fact that the CMTS downstreams come with four to eight upstreams attached to them. It is also due to the fact that the CMTS is a more complex piece of equipment, and the real estate inside the chassis is an expensive place to locate an RF power amplifier.

Another key motivation is to develop a CMTS architecture in which the number of upstream and downstream RF channels can be independently chosen and configured. Today, adding a second downstream to a DOCSIS group typically means that the operator also has to add an extra four to eight upstreams as well—something that is usually not very practical.

This flexibility which will allow CMTSs to increase their downstream bandwidth and lower the cost of the downstream, is required in order for CMTSs to more effectively compete with new high speed competitive services such as newer digital subscriber line (DSL) standards such as the 44 Mbps bonded family of asymmetric digital subscriber line (ADSL) standards—called "ADSL2"—or fiber to the home (FTTH).

Figure 1 on the next page shows a multiservice architecture where the edge QAM device is shared between a DOCSIS network and a VOD network.

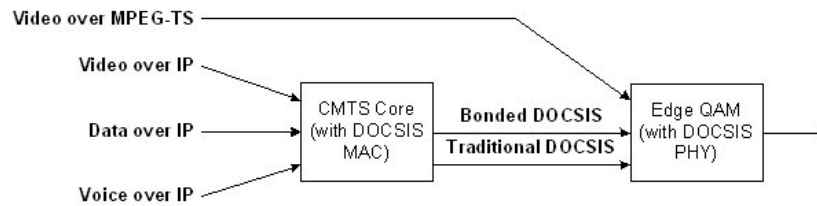


Figure 1: Multi-Service Architecture

The Edge QAM device manages two distinct types of traffic. The first is MPEG 2 (Moving Pictures Group 2) over MPEG-TS (MPEG Transport Stream) video traffic. The second is IP over DOCSIS over MPEG-TS traffic. Typically, an individual QAM channel is only carrying one of these types of traffic, although the edge QAM itself may have QAM channels that are part of either service.

The DOCSIS component provides the “triple play services” of data, voice over IP (VoIP) and video over IP. These services are combined in the CMTS core and transported to the edge QAM device typically over a switched Gigabit Ethernet network.

The data capacity of the DOCSIS downstream can be 40 to 50 Mbps for traditional DOCSIS channels with a 256 QAM downstream on a 6 or 8 MHz RF channel. The data capacity of the DOCSIS downstream can also be on the order of 200 Mbps to 1 Gbps, based upon a new emerging technology referred to in industry as “bonding” or the “wideband protocol for a

DOCSIS network” [1].

This new bonding technology is the highest priority for the upcoming DOCSIS 3.0 working groups to address.

REFERENCE ARCHITECTURE

The reference architecture for a M-CMTS system is shown in Figure 2. This architecture contains several pieces of equipment, along with interfaces between those pieces of equipment. This section briefly introduces each device and interface. Subsequent sections will go into more detail.

The edge QAM device, or EQAM for short, has its origins in the VOD environment. It is a chassis that typically has one or more Gigabit Ethernet coming in and multiple QAM modulators and RF upconverters on the output. This EQAM is being adapted for use in a modular CMTS environment.

The outputs of these devices are often referenced as just a “QAM”, rather than the full “QAM Modulator and RF Upconverter”.

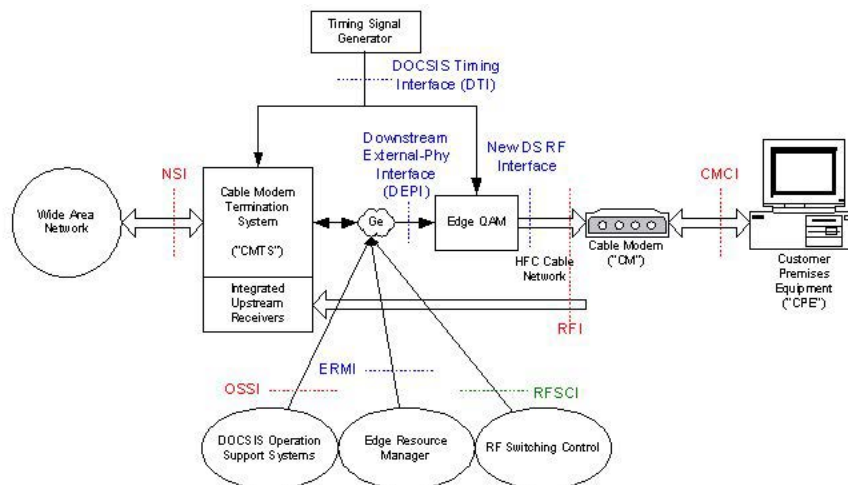


Figure 2: Modular CMTS Reference Architecture

It may be slang, but it has stuck. This paper will use the expression “QAM channel”.

The CMTS core contains everything a traditional CMTS does, except for the downstream PHY. Specifically, the CMTS core contains the downstream media access control (MAC) and all the initialization and operational DOCSIS related software.

Locating the MAC inside the CMTS core has several advantages. First, it permits the maximum reuse of existing DOCSIS software code bases. Second, it provides a local piece of hardware that can rate shape the aggregate flow of all packets to each QAM so that the output queues in the EQAM will not overflow.

This diagram currently shows the RF section of the DOCSIS upstreams internal to the CMTS core. This has been done because at this time, the Modular CMTS specifications only require the use of external downstream QAM channels, mainly because those QAM channels are leveraged off of existing equipment. External upstream QAM channels do not exist. However, there is nothing preventing an implementation of a Modular CMTS from using external QAM channels in the upstream.

The Timing Signal Generator provides a common frequency of 10.24 MHz and a DOCSIS timestamp to all MACs and PHYs.

The Downstream External PHY Interface (DEPI), is the interface between the CMTS core and the edge QAM. More specifically, it is an IP tunnel between the MAC and PHY in a Modular CMTS system which contains both a data path for DOCSIS frames and a control path for setting up, maintaining, and tearing down sessions.

If there was an external QAM burst demodulator device for the upstream

direction, then extrapolation of the name DEPI to the upstream results in upstream external PHY interface, or UEPI. UEPI is not currently being defined in the CableLabs' committees, so its definition currently is proprietary.

Downstream Radio Frequency Interface (DRFI) is intended to capture all the current and future RF requirements for the downstream direction for both integrated DOCSIS CMTS systems, modular DOCSIS CMTS systems, and VOD EQAM systems. This is quite the undertaking!

One of the goals of DRFI is to be able to specify a simultaneous system with up to 119 digital carriers and the remainder; analog carriers. Another goal is to define a high density RF connector as it is getting too difficult to use the standard F connector on high density front panels. The specification is also looking at new modulation error rate (MER) and out-of-band (OOB) noise specifications.

DOCSIS Timing Interface (DTI) is an interesting interface. DTI is a point-to-point interface from the Timing Signal Generator to all MAC and PHY components. DTI has the concept of a DTI server and a DTI client. The DTI server is the Timing Signal Generator, while each MAC and PHY has a DTI client. The DTI client is light weight. The DTI server distributes a 10.24 MHz frequency over unshielded twisted pair (UTP) with a timestamp modulated on it. The DTI returns a copy of the timestamp. The DTI server can then measure the difference in the transmitted and received timestamps to measure the round trip delay to each element. It then adjusts the transmitted timestamp to each network element so that they will all have the same sense of time.

Edge Resource Manager Interface (ERMI) is an interface that permits integration with a next generation VoD network. ERMI is used

to interface to a Resource Manager. This resource manager then allocates QAM resources to either VOD or DOCSIS applications. ERMI, however, does not directly manage within the DOCSIS QAM, so it is not a required interface on DOCSIS only systems.

Radio Frequency Switching Control Interface (RFSCI), is intended to manage an RF switch which would be used for redundancy. This would allow, for example, a bank of upstream or downstream “working” QAMs to be physically swapped out with a bank of “protect” QAMs in the event there was a failure in the “working” QAMs. This interface is not being defined as part of the first round of definitions at CableLabs.

Operations Support System Interface (OSSI) provides the management interface to each system component. One of the interesting tasks the OSSI system has to define is which DOCSIS device initializes the PHY level parameters of the QAMs such as modulation. Should the EQAM do this since it already has to do this for VOD QAMs? Should the CMTS core do this so that OSSI and CLI structures for the CMTS can be similar to what they are today? Whatever the decision, one device should get configured, and the other device should learn these values over DEPI.

The Network Side Interface (NSI) is unchanged from some nine years ago. It is the physical interface the CMTS uses to connect to the backbone network. This is typically 100 Mbps or 1 Gbps Ethernet.

Cable Modem to Customer Premise Equipment Interface (CMCI) is also unchanged and is typically 100 Mbps Ethernet.

DEPI OPERATION

DEPI is an IP tunnel that exists between the DOCSIS MAC in the CMTS core and the DOCSIS PHY that exists in the edge QAM. DEPI's job is to take either formatted DOCSIS frames or MPEG packets and transport them through a Layer 2 or Layer 3 network unharmed.

The base protocol that DEPI has chosen is the Layer 2 Tunneling Protocol Version 3, or L2TPv3 for short [2]. L2TPv3 is an Internet Engineering Task Force (IETF) protocol that is a generic protocol for creating a “psuedowire”. A psuedowire is when a Layer 2 protocol is passed transparently over a Layer 3 network. Examples of protocols supported by L2TPv3 include asynchronous transfer mode (ATM), High-level Data Link Control (HDLC), Ethernet, Frame Relay, and Point to Point Protocol (PPP).

Figure 3 on the next page shows the format of an L2TPv3 packet. There is a bit called the “T” bit in the header of each packet to distinguish between data and control packets. There is also a 32 bit session ID. The UDP header is optional with L2TPv3, but is included to permit addressing of QAM channels with the UDP destination port.

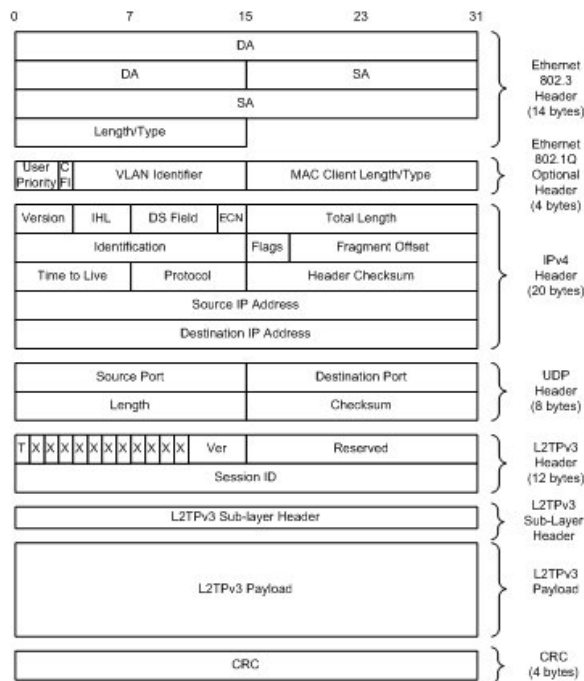


Figure 3: L2TPv3 Data Packet Format

L2TPv3 then permits a sub-header to exist whose definition is specific to the payload being carried. There is a specific sub-header for DOCSIS and MPEG-TS that currently are under definition.

The control channel allows for signaling messages to be sent back and forth between the MAC and PHY. Typical control messages will set up a “control connection” between the MAC and PHY, and then set up multiple sessions. Each session can have a different DiffServ Code Point (DSCP), and/or support a different encapsulation protocol.

There are two basic encapsulation techniques under consideration for DOCSIS. The first is a something called the Packet Streaming Protocol (PSP) and allows DOCSIS frames to be both concatenated to increase network performance, and fragmented, in case the tunneled packets exceed the network MTU size. This encapsulation is intended to carry traditional DOCSIS frames.

The second encapsulation is straight multiple: 188 byte MPEG-TS packets are placed into the L2TPv3 payload with a unique sub-header which contains a sequence number so packet drops can be detected. This encapsulation is intended to carry MPEG-TS based bonding or “traditional DOCSIS” where it is not necessary to send MAPs separately.

One of the technical considerations of the Modular CMTS architecture is its impact on the round trip REQ-GNT delay time. This is the time from when a CM launches an uncontended REQ to when it receives a MAP message with the GNT opportunity in it.

To prevent the MAP from being slowed down by data traffic, the MAP may be sent in an independent L2TPv3 session that has a unique DSCP. This DSCP will have a “per hop behavior (PHB)” that will give MAPs the highest priority and lowest latency service.

EDGE QAM OPERATION

Figure 4 shows a high level block diagram of an edge QAM that is capable of handling either video MPEG traffic or DOCSIS traffic. The expression “MPT” is an acronym for MPEG Transport.

The next interface supported is DOCSIS MPT. This is similar to Transparent MPT except that edge QAM must search for DOCSIS SYNC messages and correct them based upon the edge QAMs internal timestamp which has been derived from DTI. This mode is intended for DOCSIS frames

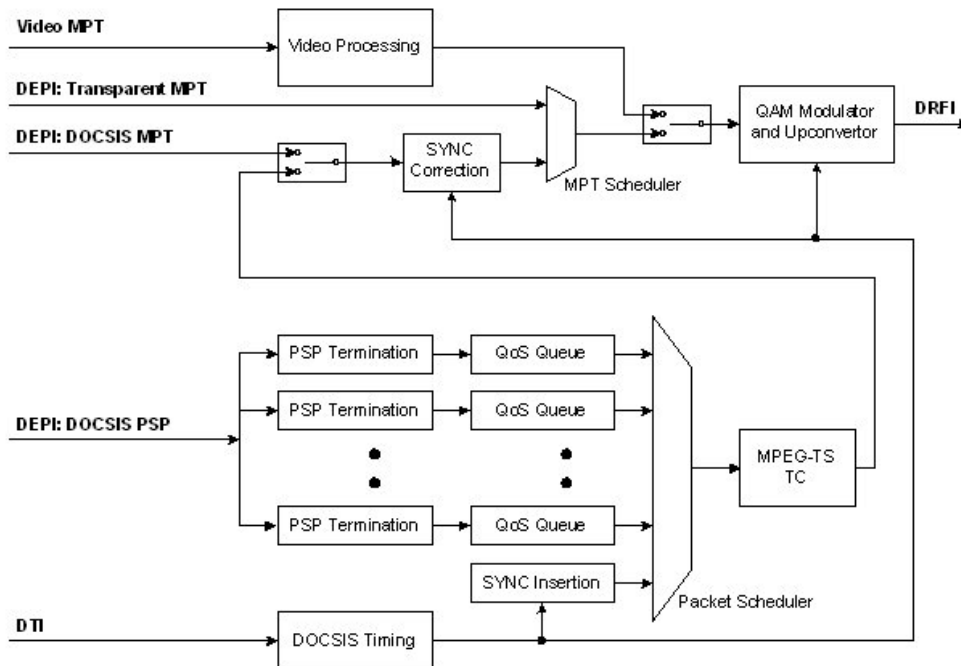


Figure 4: E QAM Block Diagram

The first interface that is shown is the VOD transport. VOD SPTS or MPTS streams are received with a format of MPT over UDP. The video processing functions generally include de-jittering, PID remapping, signaling insertion, and PCR timestamp correction. These functions are not defined in the current round of CableLabs specifications.

The next set of interfaces are the DEPI interfaces. The first one is called “Transparent MPT”. This is a simple mode in where the incoming MPT frames are copied directly to the QAM channel without any interpretation or modification. This mode is intended for MPT based bonding algorithms such as described in [1].

where the MAP is embedded into the stream and network latency is not a concern. This mode is seen as a transitional mode which may be more of interest to early implementations.

The next interface is DOCSIS PSP. Here DOCSIS frames and MAPs are received on different sessions. The DOCSIS frames may have been “streamed together” by PSP into a uniform byte stream. The PSP reassembly engine removes this overhead and recovers the DOCSIS frames. The PSP scheduler then allows MAPs to be placed in order, ahead of data and SYNC messages to be inserted. The output is then fed to a transmission convergence layer which converts the results to a DOCSIS MPT stream.

The last interface is the DTI interface which provides a common frequency and timestamp. The reference frequency is used to synchronize the downstream baud rate for use with DOCSIS 2.0 Synchronous Code Division Multiple Access (SCDMA) cable modems. The timestamp is used for the DOCSIS SYNC correction.

SUMMARY

A new CMTS architecture called the Modular CMTS architecture has been described. This architecture has the DOCSIS MAC and PHY split into two network elements. This allows each network element to be optimized both for performance and cost.

This new Modular CMTS will also provide the foundation for an entirely new class of CMTS which will have much higher data capacities than anything out there today. These new machines will require multiple 1 gigabit or 10 gigabit backhauls rather than the 100 baseT backhauls in use with today's CMTSs.

REFERENCES

1. Chapman, John T., "*The Wideband Protocol for a DOCSIS Network*", Proceedings of the SCTE Emerging Technologies Conference, January 2005
2. Townsley, Mark et. al, RFC 3931, "*Layer Two Tunneling Protocol - Version 3 (L2TPv3)*", IETF, February 2005

ABOUT THE AUTHOR

John T. Chapman is currently a Distinguished Engineer and the Chief Architect for the Cable Business Unit at Cisco Systems in San Jose, California. As a founding member of the Cisco Cable BU, John has made significant contributions to Cisco and the cable industry through his pioneering work in DOCSIS and development of key technologies and concepts critical to the deployment of IP services over HFC plants.

Included in these achievements are being the primary author of significant portions of the DOCSIS and PacketCable specifications as well as the originator of DOCSIS Set-top Gateway (DSG) and evolving specifications for DOCSIS Wideband and Modular CMTS architectures for the industry's Next Generation Network Architecture (NGNA) initiative. John has also published a number of ground breaking whitepapers on Multimedia Traffic Engineering (MMTE), DSG, QoS, and high availability and is a respected and frequently requested speaker at industry events.

John has 18 patents issued and 27 patents pending in a variety of technologies including telephony, VoIP, wide area networking, and broadband access for HFC cable networks. In his spare time, John enjoys spending time with his wife and two daughters. John is a 6th Degree Black Belt Master in Tae Kwon Do and enjoys white water canoeing and skiing.

Previous papers by John may be found at <http://www.johntchapman.com>

TRANSPORT, CONTENT, AND SERVICE IMPLICATIONS ON VOD NETWORK TOPOLOGY

George Kajos, Vice President of Engineering
Conrad Clemson, Sr. VP, Technical Operations
Broadbus Technologies

Abstract

Video on Demand (VOD) is evolving and growing rapidly. As a result, transport, content, and service offerings are changing the fundamental economics and operational efficiency of Video on Demand networks. On Demand services of a few premium movies and HBO On Demand are giving way to the NFL on Demand, Nickelodeon on Demand, Everything on Demand. Content distribution networks comprised of a small array of catchers directly connected to the back end of a small video server farm are evolving into full fledged propagation services and hierarchical storage models. The streaming network has evolved from a network with streaming servers directly connected to ASI ports to a GbE based transport network. Services are expanding from traditional transactional, free, and subscription services to include a new variety of On Demand service offerings. Everything from reality TV, to advertising, to personal ads is becoming available on line. With each of the elements of the On Demand system dissected into its pieces, the paper will put these elements together in a single cohesive view of optimal On Demand network topologies based on the evolution of transport, content and service type.

OVERVIEW

The VOD environment is clearly evolving in multiple simultaneous dimensions. Within this environment, there are several competing architectures for appropriate VOD server

deployments. Some architectures propose decentralized VOD server deployments; others propose centralized server deployments, finally some compromise with a hybrid approach.

Initially, content usage data from a small VOD installation is examined. Then, this paper evaluates the effect of content placement on the transport network and the economics of a complete VOD solution as a function of centralizing vs. decentralizing servers. In the face of dramatically decreasing transport costs, the paper identifies the few scenarios in which edge caching may be an effective approach to certain VOD applications. The paper also examines content propagation and replication. Finally, this paper examines how the evolution of new services, some of which may be personalized to a per subscriber basis, should dictate the placement of both the servers and the content.

SERVICES

The last half decade has validated Video on Demand. Hundreds of VOD deployments have occurred featuring MOD, SVOD, and FOD [CED]. A consistent server design point had previously been 500 – 1,000 streams and 1,000 – 2,000 hours of content. Server clustering allowed installations to support 10,000 – 20,000 streams.

Accordingly, these successes have unleashed the potential demand for a wide range of new services. Examples of new services being trialed or conceived include:

- Music on Demand
- Non-linear Live Broadcast
- Network PVR
- Customized content
- HD Content and Widescreen format for all the of above

It is assumed that Music on Demand will not have a major impact on streaming and storage capacity.

HD content is still evolving. The impact of the current format is a 4X multiplier on streaming bandwidth and storage requirements as MPEG 2 HD content is being transmitted at 15 Mbps. The demand for HD is steadily increasing.

MOD Wide Screen versions appear as just another content. They are typically equivalent in size.

Network PVR and Live Broadcast services have the potential to greatly impact capacity requirements. Quantifying the required bandwidth is straightforward. Service offerings and business rules determine temporary storage requirements. For example, a service offering of 100 SD and 10 HD channels has an ingest requirement in excess of 500 Mbps and a content storage requirement, temporary or permanent, of 250 Gbytes/hour. The selection of content to retain and the duration it is to be made available could vary widely based on the business rules of the offering.

SERVER DESIGN

The continuing advance of technology allows this generation of video servers to break the dependency of streaming capacity on disk bandwidth. First generation video servers relied on streaming from hard drives. Consequently, the number of streams served

was a one-to-one relationship with the bandwidth available from hard disks. Extremely sophisticated striping and scheduling techniques were employed to drive up stream counts. Moreover, custom trick files were prepared for fast forward and rewind in order to remove the variance from disk access. In general, this meant a trick file for every fast forward/rewind rate or a single fast forward/rewind rate.

A hard drive based server also posed limitations on the amount of content which could be ingested. Updating hard drives with new content reduced streaming bandwidth while disk writes were scheduled. In an environment with just MOD and SVOD services, content propagation could be scheduled at off hours with little adverse affect. However, as we explore live broadcasts and real time propagation, the required inbound content loading bandwidth goes up considerably.

This generation of video servers is designed around two principles:

- Independent scaling of streaming bandwidth and storage capacity
- Real time content ingest and turnaround performance

Server architecture designers must carefully consider the tradeoffs between the following:

- Processor performance
- Disk I/O bandwidth
- Network bandwidth
- Memory bandwidth
- Backplane/Interconnect I/O bandwidth

SERVER PLACEMENT AND NETWORK DESIGN

Extensive research has been conducted and numerous papers have been written on the topic of video server network topology. Most authors describe the approaches as centralized, decentralized, and hybrids of the two. Looking at the next-generation of server capacity and content library sizes, and taking into account both cable HFC and IP video, it is instructive to generalize the following components:

- Video Server Complex – server technology capable of streaming, local content storage, On Demand and scheduled ingest, and session/stream management.
- Transport network – the network supporting streaming to On Demand clients.
- Propagation network – the network supporting On Demand and scheduled propagation to servers.
- Content Storage Server – a generalized library and central repository for the content made available for On Demand services.
- Preparation Server – receivers of live broadcasts which then encoded for propagation to video servers to On Demand clients for play.
- On demand clients – the media decode and display point at the subscriber. Figures 1 and 2 depict centralized and decentralized VOD environments respectively.

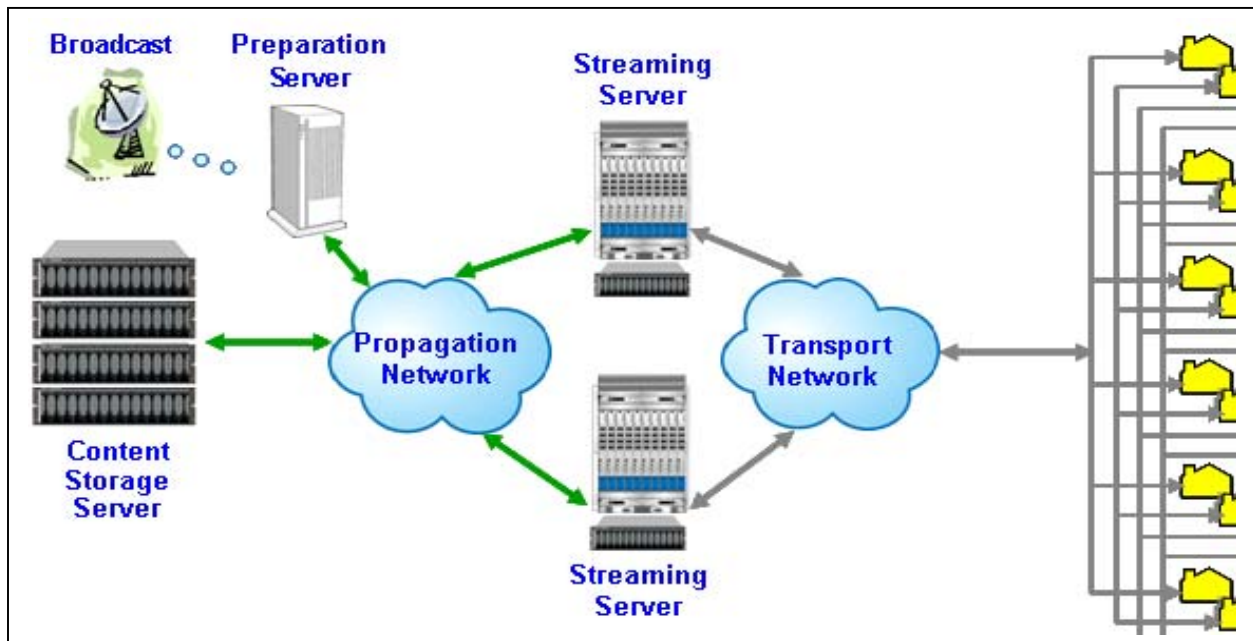


Figure 1: Centralized Server Placement

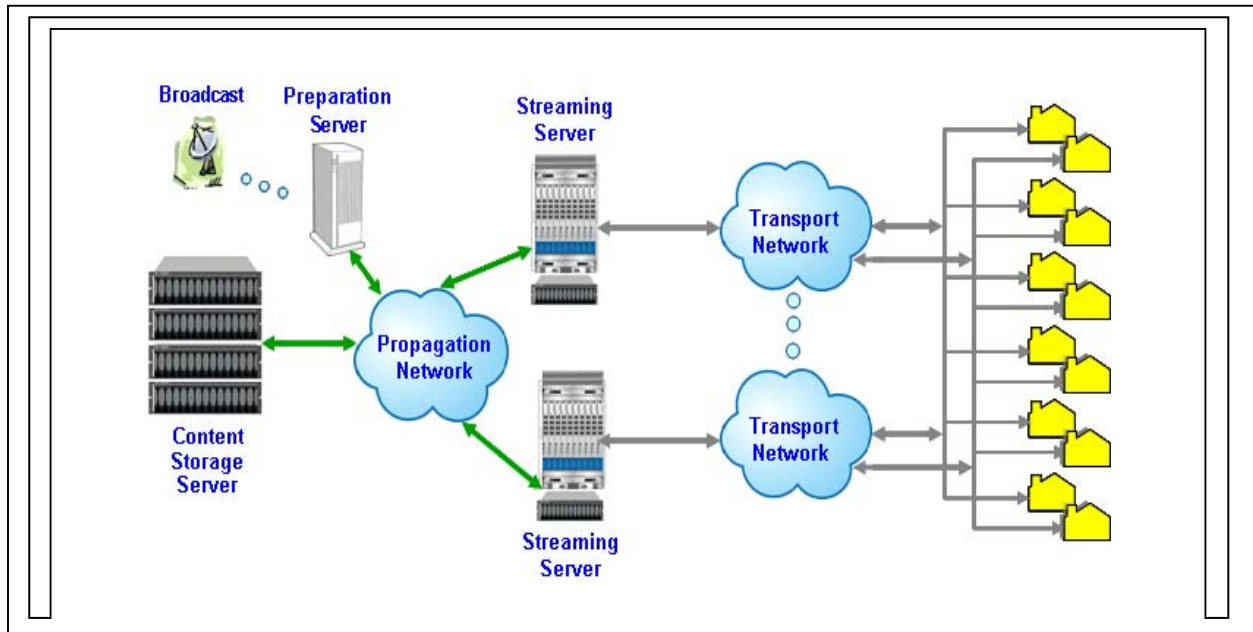


Figure 2: Decentralized Server Placement

The distinctions between these environments follow:

1. The centralized server environment has complete connectivity to the entire population of supported On Demand clients over the centralized transport network.
2. The centralized server environment receives only scheduled content and little or no On Demand content over the propagation network.
3. In the centralized environment, it is possible to load balance across the entire population of clients. In the fully decentralized environment, load balancing is restricted to the partitioned subset of clients.
4. The decentralized server environment has connectivity with a partitioned subset of the population of supported On Demand clients over regional transport networks.

5. The decentralized server environment receives both scheduled and On Demand content over the propagation network.

The differences in the two environments are chosen to emphasize the trade-off between the cost of transport network bandwidth and content replication.

The network bandwidth/storage trade-off is not the only consideration between centralized and decentralized approaches. Other considerations include:

- Replication of control components at decentralized sites.
- Operational costs of additional decentralized sites.

In reality, most VOD system designs are hybrid approaches. For example, even in a centralized environment, it is unlikely that every server need to have connectivity to

every client. Acceptable load balancing is possible with reduced connectivity. In a decentralized environment, for added reliability it would be advantageous to have more than one server capable of reaching each On Demand client. The system could operate in a degraded mode until system repair completes.

CONTENT PROPAGATION

As discussed above, in first generation video servers, content propagation requirements corresponded to the gradual refresh of new MOD, SVOD, and FOD offerings. These could be loaded onto video server complexes with little affect on server performance during low usage periods.

This generation video server must be designed for two new sources for content propagation:

- Live Broadcast
- On Demand propagation.

This paper uses the term “On Demand propagation” to refer to the requirement to move content to a server which has received a purchase request and does not have the required content. The case arises in server environments where not all content is located on every server. In a centralized server environment, a session manager could direct the request to an appropriate server.

However, in a decentralized server environment, clients are partitioned by servers. It is unreasonable to replicate all content at every site in a decentralized server environment. Consequently, this paper defines the case when content must be transferred to a server to grant a client request. Figure 3 demonstrates how On Demand propagation works. The client requests content not available on the local server or server complex. A request is made to a regional propagation server for the required content. A “filler” content is transmitted to the client at the start of the upload. The filler could consist of previews or advertising. When the server buffers enough content, play out begins.

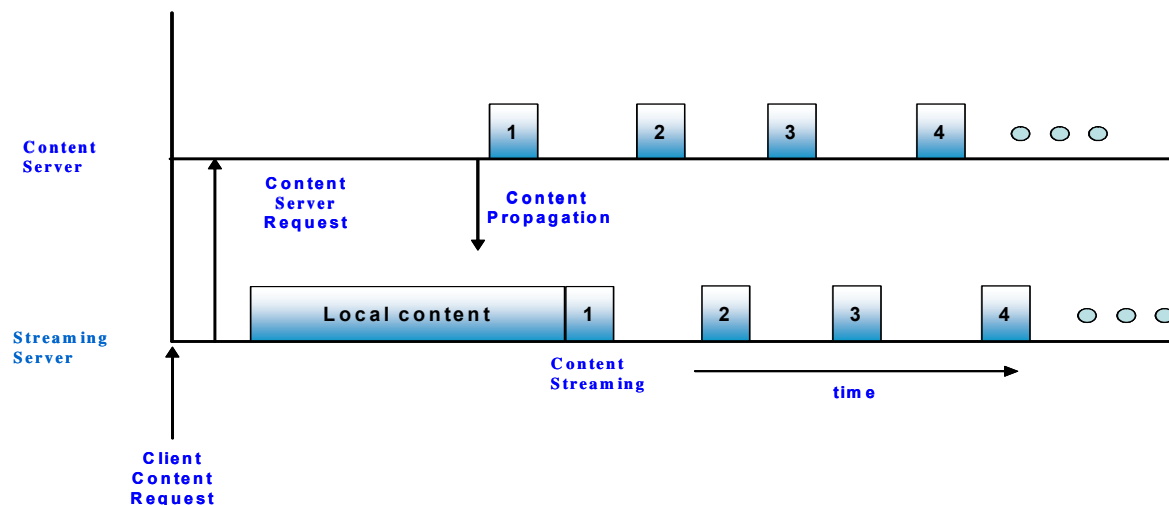


Figure 3: On Demand Content Fulfillment

Live broadcasts and On Demand propagation, move guaranteed quality of service from a server memory or disk subsystem problem to the propagation network. Unlike the dedicated transport network which is most often provisioned for the maximum stream capacity dictated by the HFC QAM capacity or as a percentage (provisioned take rate) of the total pool of clients, the design of the propagation network should include a policy on how to allocate between scheduled, live broadcast and On Demand bandwidth. As discussed earlier, Live Broadcast input bandwidth requirements could range in the 500 Mbps. Figure 4 depicts an example of how the allocation policy could vary during a 24-hour cycle.

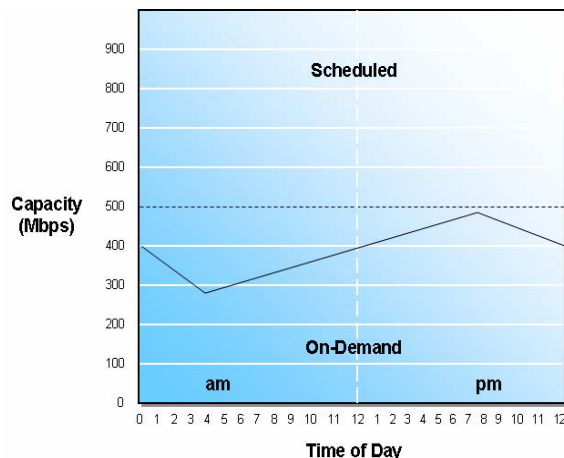


Figure 4: Network Bandwidth Allocation Policy

TRANSPORT AND PROPAGATION NETWORK TECHNOLOGY

The streaming network has evolved from a network with streaming servers directly connected to ASI ports to a GbE based transport network. Ethernet price points have continually fallen due to the ever increasing reach of Internet protocols from traditional LANs to geographically dispersed WANs. While there are

alternatives in the WAN, such as Packet over SONET and ATM, as GbE moves to 10Ge, it appears that 10 GbE will become the most cost effective, high bandwidth solution.

Consequently, this paper examines ways to interconnect geographically separated 10Ge pipes. The most straightforward are dark fibers and long reach optics to create point to point links. However, this is an inefficient use of the bandwidth available in the fiber and would only make sense if only a single trunk of 10 Gee is required.

Another approach is wave division multiplexing (WDM). WDM technology allows data from multiple sources to share a single fiber by transmitting on individual wavelengths. WDM interfaces have been incorporated in switches and multiplexers.

Two types of WDM are in use today: Dense WDM and Coarse WDM. DWDM is ideal for high bandwidth, long haul applications. Current DWDM technology can squeeze over 30 channels in C and L optical bands.

CDWM uses lower cost optics and is characterized by wide channel spacing over a wide spectrum. CWDM technology can supply 18 channels from 1270 to 1610 nm.

WDM technology is ideal for accommodating 10 Gee streaming pipes to remote QAM or DSLAM locations

OBSERVED DATA

In this section, VOD and SVOD are examined as a starting point for planning for new services. VOD today is dominated by SD (standard definition) content and is transmitted as MPEG2, 3.8Mbps/stream. It

is used for Movies on Demand and “Subscription on Demand” Services.

Figure 5 and 6 depict VOD usage for one day, January 29, 2005, at a relatively small site. The days’ totals are 5,133 streams

across 703 contents. As one would expect on a Saturday, the maximum number of sessions peaks around 9:30 PM as shown in Figure 6. The distribution shown in Figure 2 plots content usage from most used content (117 plays) to least viewed content (1 play).

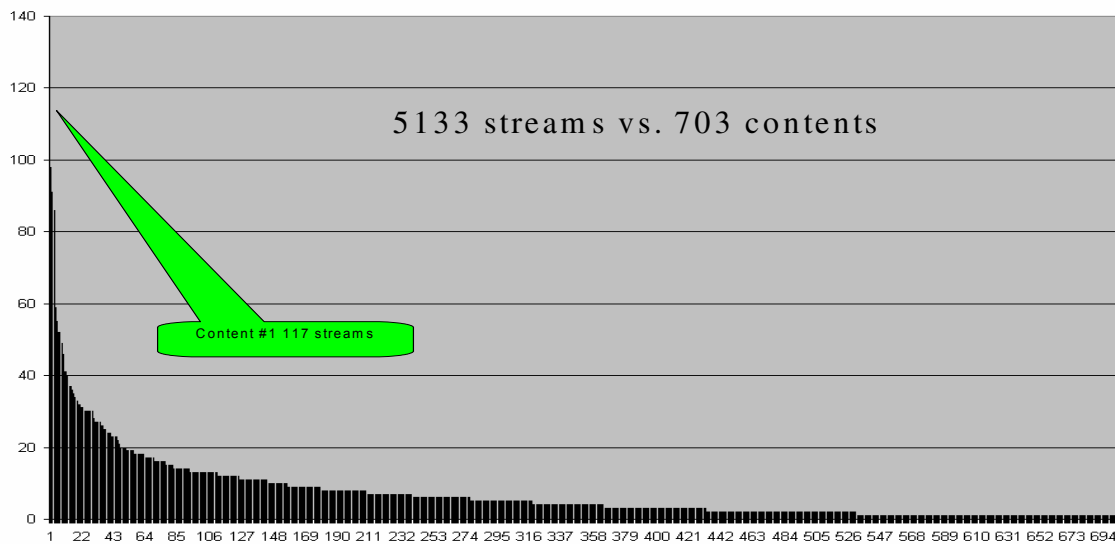


Figure 5: January 29, 2005 – Streams versus Time of Day

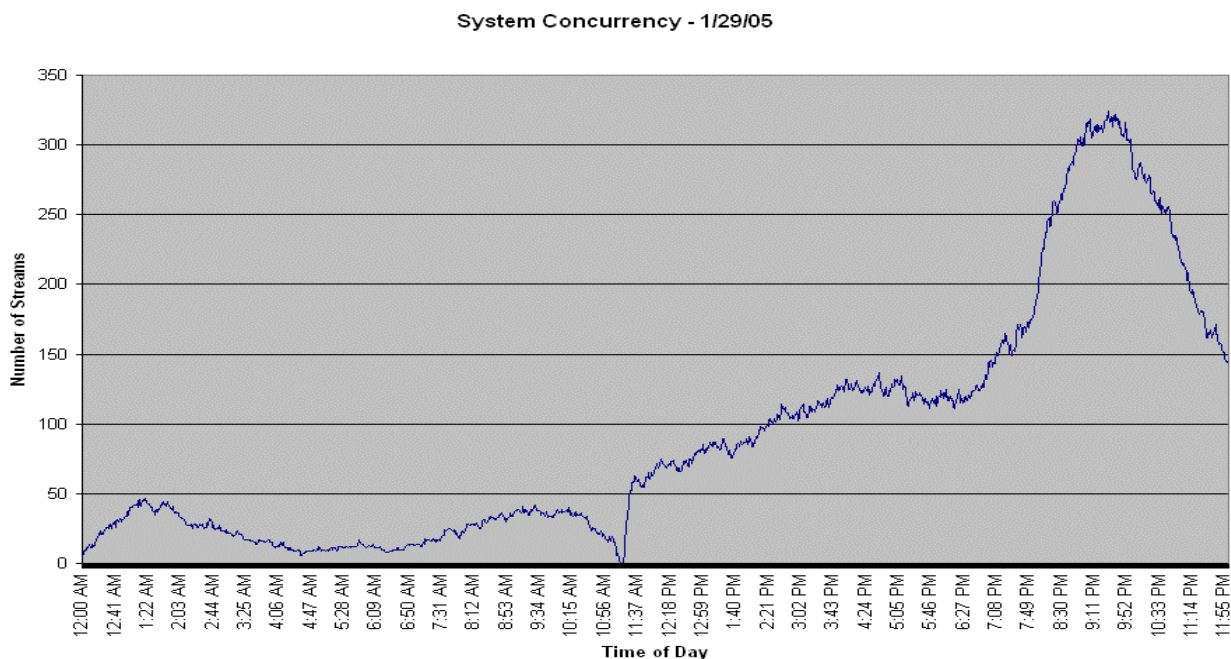


Figure 6: January 25, 2005 – Streams vs. Time of Day

These observations are typical and reported in a number of other works [5]. In

fact, the popularity distribution can be fit to Zipf’s law, which states that the probability

of requesting a program m , where $m = 1, 2, 3 \dots$ out of N movies is :

$$C/m \text{ where } C = (1 + 1/2 + 1/3 + \dots + 1/N)$$

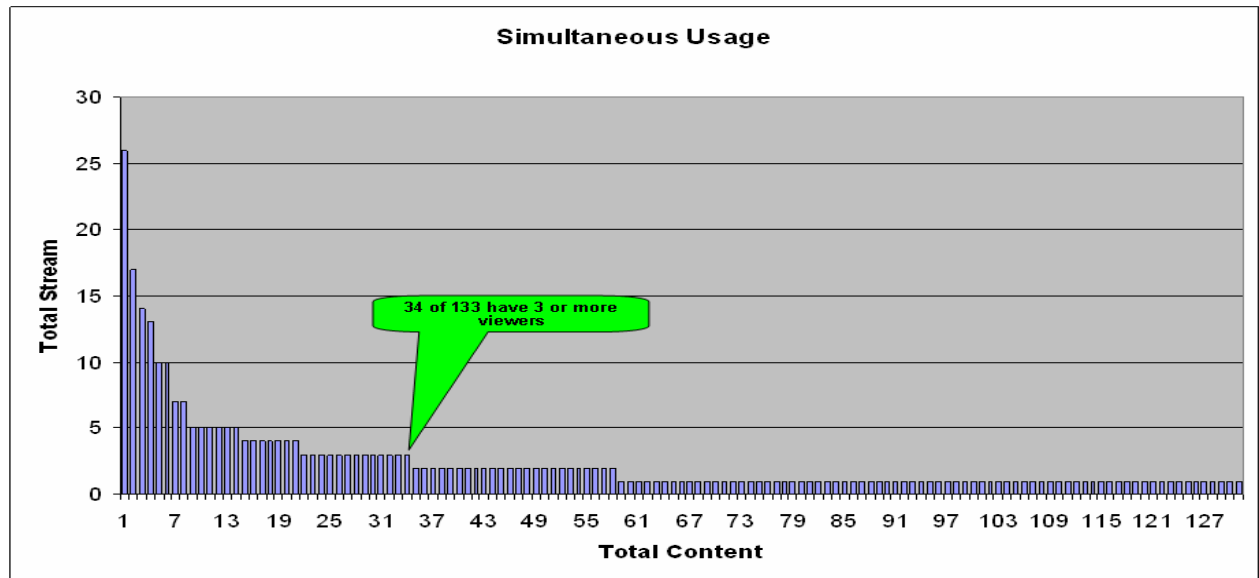


Figure 7: Simultaneous Content Viewing at Peak Usage

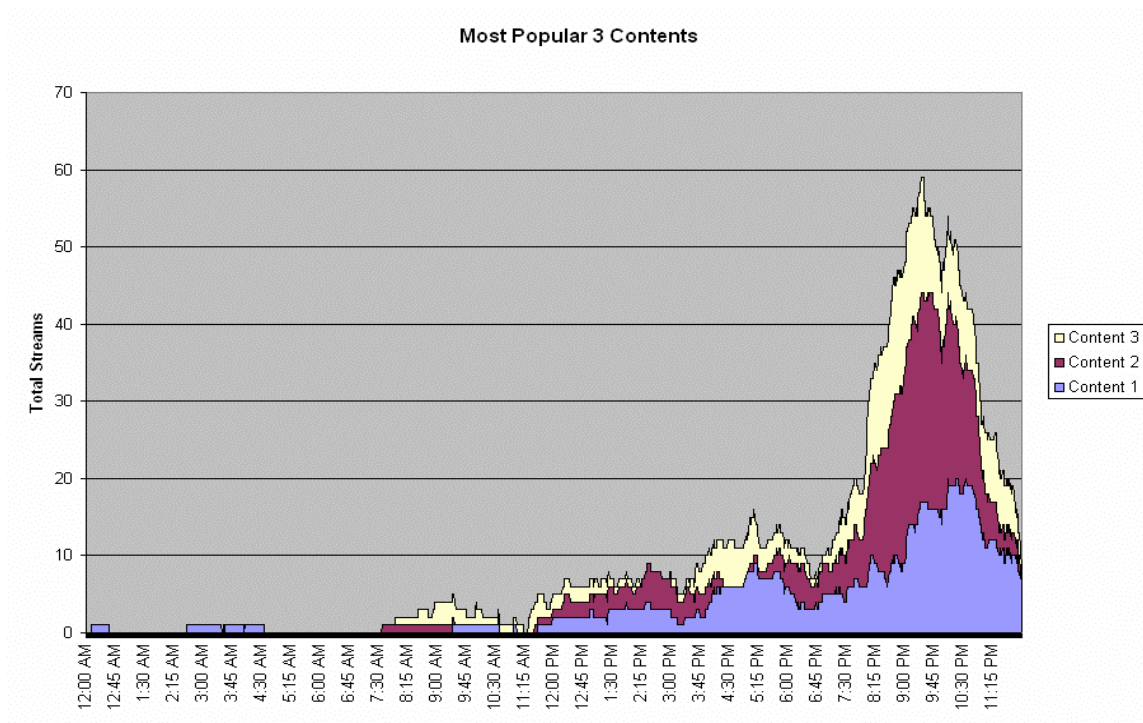


Figure 8: January 29, 2005 – Most Popular 3 Contents

Figures 7 and 8 examine the relationship between simultaneous sessions and content. In Figure 7, it can be seen that 58 of the 133

contents are being viewed by two or more set top boxes. Figure 8, depicts that at the peak of viewing, the top three contents total 57 of

the 322 sessions. As video servers are scaled, these relationships will be used to assist with the tradeoff between memory, network, and I/O bandwidth.

Finally, additionally data obtained from discussions with a number of MSO's with mature VOD deployments suggests that actual VOD usage data varies widely from service group to service group. While it is not possible at this time to publicly share specifics of these results, it is generally understood in the industry, that within a system, peak usage may vary from 2% to 14% on a service group by service group basis.

This data proposes some interesting conclusions.

First, real world data validates a Zipf distribution model across multiple deployments. The Zipf curves favor caching architectures, in general, and specifically making caching architectures more effective as the stream count grows. This is the case, because in a Zipf curve, the tail is relatively constant, while the peak of the curve grows dramatically with stream count. As centralization takes place, VOD servers that accommodate extremely high concurrency, such as caching servers, will serve more streams at a lower cost.

Additionally, the tail stays relatively constant. In hybrid architectures, it is generally assumed that less popular titles are streamed from the core to conserve the replication of storage. This data supports that model, but it also points out that a library server which has a caching capability, can perform both tasks.

Some server designers would argue that it makes sense to cache the popular titles at the

edge to conserve transport costs. While this approach has academic appeal, it does not hold up to real world scrutiny. The problem with this approach is that it assumes a uniform distribution of concurrency in each service group. The data suggests that each service group has its own peak concurrency. These server designers would propose to either provision the entire system to the average concurrency or, worse, the peak concurrency. Provisioning to the average concurrency will result in denial of service to the peak service groups and over provision the low concurrency groups. Provisioning to the peak concurrency will result in massively over provisioning the entire system.

When transport costs are equal to or less expensive than the streaming server costs, as they are today, the only logical way to provision a VOD system is to centralize the architecture. This provides the operator with tremendous economies of scale in the streaming subsystems. At the same time it allows the operator to provision across all service groups without stranding streams at the edge of the network.

NEXT GENERATION VOD ARCHITECTURE

In this section, hypothetical VOD system is created for the purpose of further exploring centralized approaches and decentralized approaches. Assume that an environment is provisioned for a take rate of 300,000 active clients. There are two server capacities available – 3,000 and 15,000 streams. In this exercise, the paper examines the number of servers and the equivalent number of gigabit Ethernet links required at each extreme. One additional consideration is the possibility that a decentralized server will need double the ingest bandwidth of a centralized server to accommodate On Demand propagation.

Table 1: Servers and Transport Links

300,000 streams	Streams per server	Servers	GbE Transport Links	10GbE Transport Equivalents	Propagation GbEs
Centralized	15,000	20	60	6	20
Decentralized	3,000	100	12		100/(200)

Some general observations can be made:

- More servers, storage, control systems, in decentralized model
- Ability to collapse ten 1GbEs into one 10GbEs in the centralized model
- Potential additional load on the propagation network in distributed model

- Higher cost transport in centralized model
- Greater potential to share storage in centralized model

Figure 9 depicts an architecture which collapses the transport network connectivity.

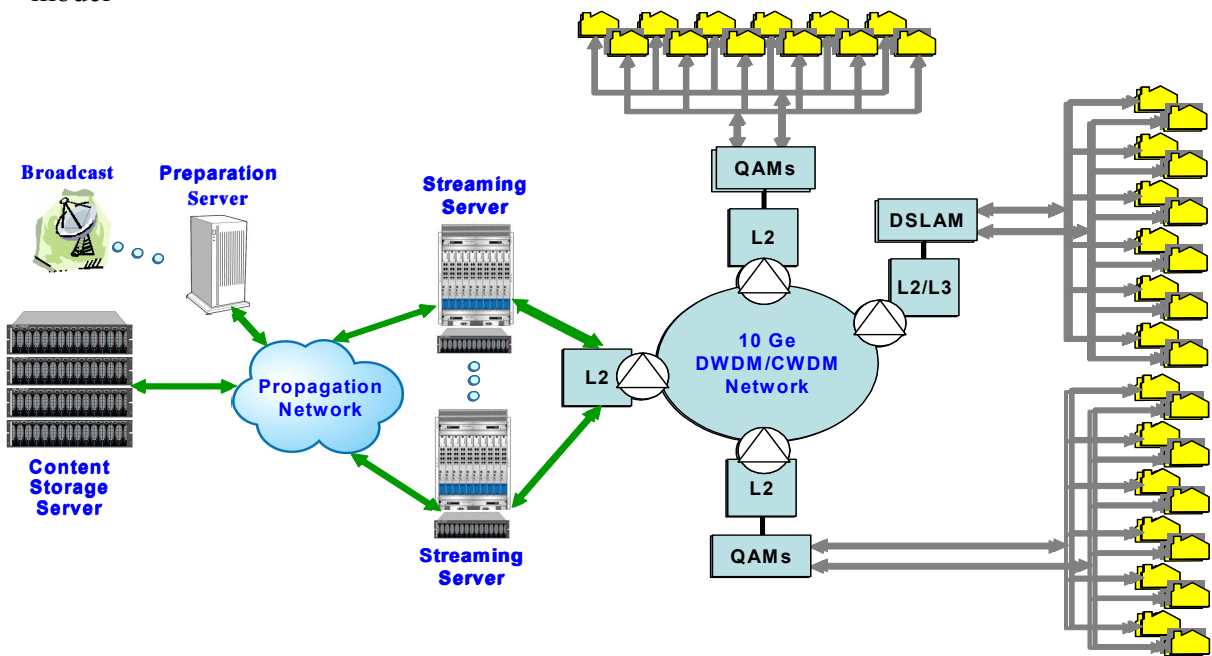


Figure 9: Next-Generation Collapsed VOD Network Architecture

The transport network is represented as a WDM optical network connected through layer 2 switches. Because 10Ge pipes are steered to specific Lambdas and partition the client space, the architecture is not a fully centralized environment. However, most of the advantages of the centralized environment are realized.

SUMMARY

During the previous half decade, many VOD installations were monolithic and self contained for ingest and streaming. Advances in server and transport technology allow new services to be considered for VOD deployments. In this paper, data regarding content popularity during a single day and at peak load was presented from one VOD installation. The next-generation VOD architecture presented in the paper is well suited to meet the scale to the requirements dictated by new services.

ACKNOWLEDGEMENT

The authors wish to acknowledge Michael Kahn for extracting and processing the January 29, 2005 data.

AUTHOR'S CONTACT INFORMATION

George Kajos
VP, Engineering
Broadbus Technologies, Inc.
80 Central Street
Boxborough, MA 01719
gakajos@broadbus.com
Office: 978.264.7906

Conrad Clemson
Sr. VP, Technical Operations
Broadbus Technologies, Inc.
80 Central Street
Boxborough, MA 01719
cclemson@broadbus.com
Office: 978.264.7905

REFERENCES

1. Skarica, Christopher, and Joe Selvage, "Environmentally Hardened optical CWDM Solution for MSO Access Network Capacity Expansion," Cable-Tec Expo, June 15-18, 2004.
2. Kajos, George, "VOD Storage Architecture for Hierarchical Network Content Libraries," Cable-Tec Expo, June 15-18, 2004.
3. Jancola, Mark, "Considerations in Scaling Video-on-Demand Systems," Cable-Tec Expo, June 15-18, 2004.
4. "North American Cable Deployments," CED Magazine, December 2004.
5. Vassiliakis, Constantine, Michael Peterakis, Peter Triantafillou, "Video Placement and Configuration of Distributed Video Servers on Cable TV Networks," Multimedia Systems, 8:92-104, 2000.

UNIFIED DATA AND VIDEO CMTS: ONE SYSTEM FOR ALL SERVICES

Scott Cummings and Victor T. Hou
Broadcom Corporation

Abstract

This paper opens with a quick history of the traditional Head End (HE) and current CMTS platforms in the HE. The “loosely coupled, tightly bolted” architecture of the traditional HE is exposed. The issues in adding new services and technologies will become apparent.

Then an overview of several current and new technologies entering the next generations of HEs is listed. Each technology is described. Then the Modular CMTS (M-CMTS) architecture is applied to the Next Generation Network Architecture (NGNA) for the HEs. The traditional CMTS is subdivided into three modules: CMTS core, Edge QAM, and Upstream Receiver. A new component to the HE is defined: a System Resource Manager (SRM).

The incoming technologies are applied to the SRM and M-CMTS. As these technologies are incrementally added to the SRM and M-CMTS, the changes to support them are simple. Adding servers, SRM interfaces, and possibly CMTS cores and the system is ready for the new technology or service. The combination of the SRM and M-CMTS convert this HE architecture into the unified CMTS, a system for all services.

Traditional Head End (HE)

The traditional HE is composed of many devices that offer several services. Video, Data, and Voice services are all supported by the HE. Each of these services and the equipment used to deliver these services will

be quickly examined. The term HE is used a bit loosely here. There are distribution hubs and other types of nodal locations that may provide the functions described in this paper as a HE.

Cable Modem Termination System (CMTS)

The current DOCSIS 1.x, and 2.0 CMTS platforms predominantly deliver data service for cable modems. These systems operate with a downstream modulation type of either 64-QAM or 256-QAM in the downstream. The downstream channels support either 6 MHz (North American plants) or 8 MHz (European plants). The CMTS also provides some networking services for the cable modems (DHCP, ToD, Gateway). Some CMTSs incorporate these servers, others do not. DOCSIS 1.1 introduced Quality of Service features to the CMTS such as rate shaping and per user rate limiting.

The CMTS also provides upstream data services. DOCSIS 1.1 is limited to a maximum bandwidth for a single channel of 10 Mbps. To provide symmetric bandwidth with the downstream, the DOCSIS 1.1 CMTS platform are designed for a 1:N downstream to upstream ratio, where N would be 4, 6, or 8. DOCSIS 2.0 introduces many upstream modulation types which now allows the CMTS to offer QPSK, 8-QAM, 16-QAM, 32-QAM, 64-QAM, or 128-QAM (SCDMA only). However, the downstream to upstream ratio has not changed.

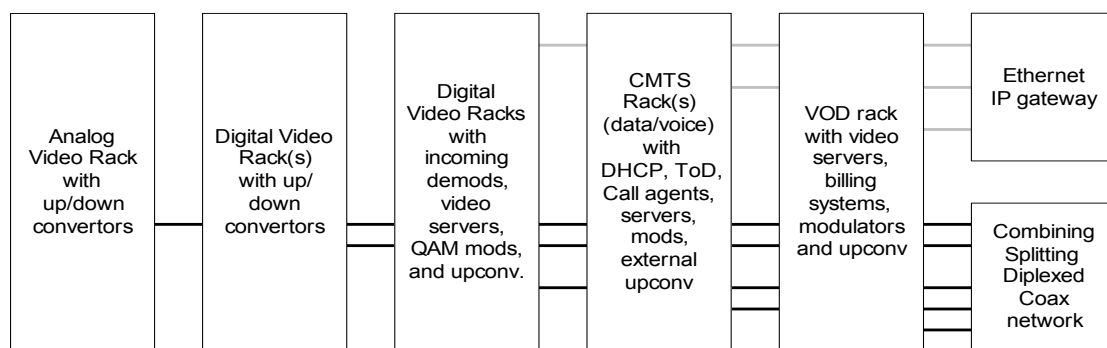
The current CMTS platforms are composed of a rack of blades. Typically the CMTS contains a network blade, a processor blade and a physical layer blade. Additional

blades can be added to the rack to meet the needs of the node or head end supported. The physical layer blade maintains the 1:N ratio. These are standalone systems that merely share the Hybrid Fiber Coaxial (HFC) cable plant.

The current CMTS platforms also delivers Voice over Internet Protocol (VoIP) service. The VoIP services are bidirectional services that have restrict latency and jitter requirements relative to the general data services. Due to these requirements, the VoIP services are often segregated from the data

combination of both analog and digital signals are transmitted on the same HFC plant.

The second type of video is narrowcast video. The narrowcast video is composed of Video on Demand (VoD) and other selectable video services. These services are not streaming video all the time. A given video stream may be streamed on multiple RF channels. Billing systems for the individual purchase of video streams are also required for these services. The narrowcast video services often also require a return path for ordering and billing confirmation.



Traditional HE system level diagram

services. This segregation may be as simple as dedicated channels for VoIP, or as rigorous as a dedicated CMTS for the VoIP services.

The CMTS systems make up only a small part of the Traditional HE.

Video Distribution

The majority of equipment in a traditional HE is for video distribution. There are two distinct types of video distribution in a HE.

The first is broadcast video. The broadcast video may be delivered with traditional NTSC analog signals, or with digital video MPEG transmissions. In most HE today, a

Narrowcast video may also be used to send programs only when they are viewed. Rather than broadcast all the channels available to the subscriber, a future model for narrowcast video would send only the channel requested in either multicast or unicast streams.

Ignoring the equipment just to receive video (multiple satellite dishes, local antennas, high speed fiber connections, downconverters), the HE is composed of multiple QAMs, upconverters and MPEG multiplexing stations. The MPEG multiplexing stations may be simple multiplexing units or very sophisticated and expensive units that are capable of transcoding and transrating MPEG streams in real time.

The distribution of broadcast video and narrowcast video both use the same type of equipment. The billing, switching, and return path requirements of the narrowcast video generally create implementations for these systems that are stand alone systems relative to the broadcast video systems.

The CMTS systems providing data and VoIP services also tend to be independent systems from the video distribution systems. This creates a collage of equipment in the HE. This also creates operational issues for maintaining a HE. Each system requires individual expertise to be maintained and operated. As new technologies and/or new solutions for existing services become available, these systems tend to contain requirements that lead to stand-alone implementations as well.

New Technologies

There are several new technologies that are currently available or are soon to be available to the next generation of HEs. These technologies include DOCSIS Set-top Gateway (DSG), Internet Protocol (IP) video, Channel Bonding, the Personal Video Recorder (PVR), MPEG-4, part 10 (AVC) digital video compression and 1024-QAM downstream modulation. Each of these innovations has targeted applications and benefits.

DSG

The DSG technology provides a DOCSIS tunnel to any device connected to the HFC plant. This tunnel is a 1 way (downstream only). However if the CPE device contains some form of an embedded CM (eCM), a 2 way communication path can be implemented. (both downstream and upstream). This opens a development path for the HE to control devices other than Cable Modems (CMs) with DOCSIS. DOCSIS has provided the MSO

with the certification process to guarantee compatibility across multiple vendor devices. The certification and qualification process also delivers a very robust and reliable system. The Consumer Premise Equipment (CPE) connected to the HFC plant can now all be controlled using this DSG tunnel.

The DSG technology is available today. The DSG effort has been active at CableLabs and the specifications are complete.

IP video

The concept of delivering video over IP on an HFC plant has been around for quite some time. IP video dominates video delivery to Personal Computers (PCs). QoS, buffering issues, and overhead issues have all hampered the acceptance of IP video. IP version 6 (IPv6) has addressed some of the QoS issues. Memory and disk space technology continue to follow Moore's law. The cost for a bit of RAM or disk space continues to drop. This has greatly reduced the buffering issues with IP video, and will continually reduce buffering issues well into the future. The additional overhead for the IP headers can be addressed with Payload Header Suppression (PHS) in DOCSIS CMTS platforms.

The power of IP video comes from a couple of sources: economies of scale, and the robust feature set of the IP protocol. The IP technologies are ubiquitous. These solutions are extremely cost effective from an equipment standpoint. IP video also automatically provides several avenues for advanced video services. These services include stopping, pausing, forwarding, rewinding, book marking, and other video control features. The IP video can be served to a wide range of devices. PCs, STBs, and even Hand Held Devices (HHDs) are all capable of processing IP video.

PVR

The PVR has been available to the consumer for over a year now. The original PVRs provided storage to hold recorded video programming and a simple software Graphical User Interface (GUI) to help program and manage stored recordings. The early PVRs had limited storage and therefore did not offer many features beyond a mechanism to record and store video.

This is similar to the MP3 player. The MP3 player started as a simple device to store and record audio. Over time, the MP3 player has expanded to be a general purpose storage device that has the capability of playing and recording audio.

The PVR has the opportunity to take a similar migration from a dedicated video device to at general purpose storage devices that can play and record video, audio, still photos, and even display web pages. The PVR actually could be the link between the STB and the PC. The STB has had a processor, memory, an Operating System (OS), and a GUI for several years. Many of the later models of STB have Ethernet and/or USB ports on them. Keyboard style remote control units exist for both STB and high end TVs. The lone missing piece of a PC in an STB is the disk drive. The convergence of the Web and Video, the STB and the PC, could all be in the PVR.

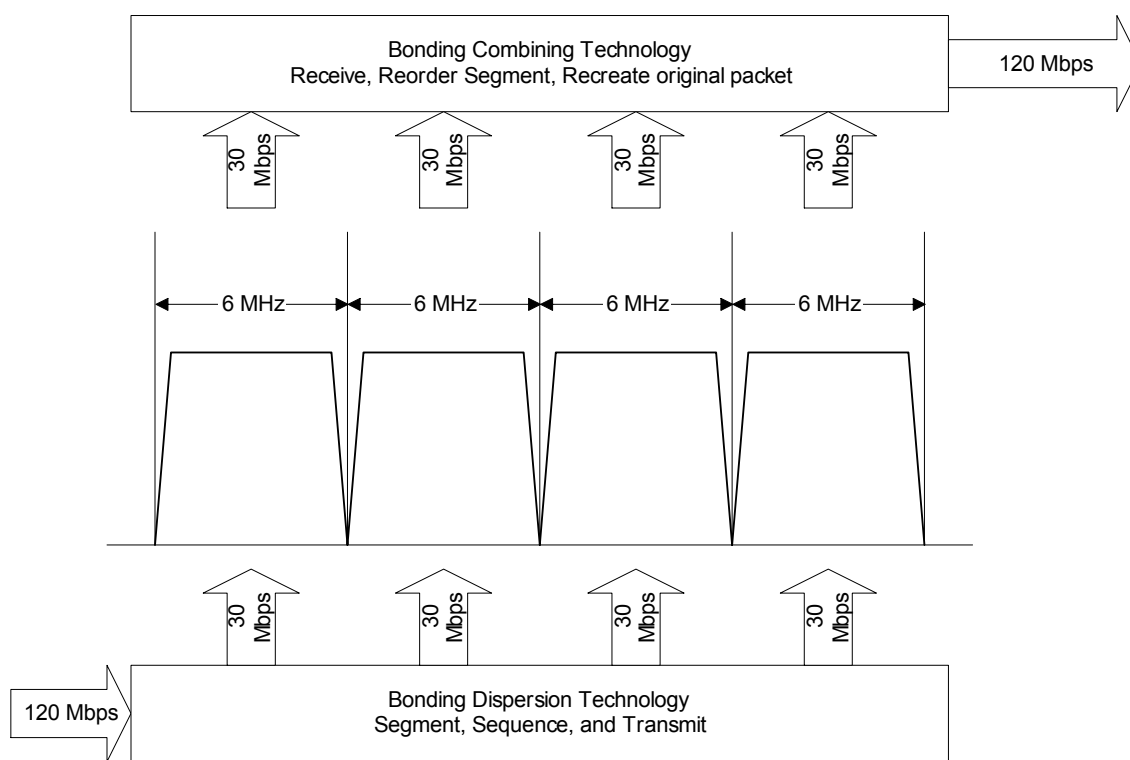


Diagram of channel bonding a 120 Mbps load across 4 downstream channels

Channel Bonding

Channel Bonding is a new technology that is part of the discussions for DOCSIS 3.0. The concept of channel bonding is fairly

simple. Channel Bonding bonds N number of existing channels into 1 larger virtual channel. This provides several features. First, the maximum throughput to a single CPE will be increased by a factor on N. An MSO could

offer a high speed service of well over the 38.8 Mbps that a single 256-QAM downstream channel can provide today.

Another feature of Channel Bonding is the statistical multiplexing gain of having more users share a wider channel. Statistical multiplexing benefits occur when a channel gets oversubscribed. Service on that channel has to be reduced or stopped for some users. If another channel was undersubscribed, the excess load could be taken by the second channel. Today, instantaneous load balancing is not available in CTMS platforms. If both channels could be virtually combined, the new bonded channel would not be oversubscribed and no reduction or loss of service would be observed by the users.

The granularity of the Channel Bonding solution may also provide other benefits. For Channel Bonding to work, the CPE must have the ability to simultaneously receive multiple channels. With more channels to select where and when to place downstream packets, latency and jitter could be better controlled. This is similar to the statistical multiplexing gain with the expectation that the system has not been loaded to the point of over subscription. As the system is approaching the maximum subscription rate, this level of the loading or congestion begins to create an increase in packets that request the transmission at the same transmission time. Clearly in a single channel, two packets cannot be transmitted at the same time. One packet will be transmitted, and the other packet will either be latent or get advanced. In this case, jitter has been introduced. As the system approaches the maximum subscription rate, jitter and latency begin to increase.

A Channel Bonding solution that has MPEG layer granularity and the freedom to select the channel at transmission time can send multiple packets at the same time. If a channel is at the maximum subscription rate,

the channel bonding system can use the other channels to help get critical packets downstream with minimal jitter and latency effects.

1024-QAM modulation

Many new technologies have been applied to the HFC plant. No advancement in DOCSIS from its inception has provided improved bandwidth in the downstream. PHS can be applied in the downstream to save some bandwidth. The only improvements in bandwidth in the downstream over the past 5+ years has been the expansion of the HFC plants frequency spectrum. The original HFC plants started with an upper frequency limit of 550 MHz. This has expanded to 650, 750, 850, and now there are discussions of opening up to more than 1 GHz. The frequency expansion has hardly been a new technology. Replacing cable, connectors and amplifiers with cleaner, equipment specified up to the higher frequencies is not an architectural technology improvement.

A more efficient modulation scheme can be introduced to the HFC plant. Currently, the highest modulation scheme available to an HFC plant is 256-QAM. This provides 8 bits per symbol. 1024-QAM provides 10 bits per symbol. Upgrading the modulators and demodulators to 1024-QAM would improve the raw bandwidth by 25%.

Beyond the straight bandwidth gain of the modulation order, a statistical multiplexing gain will also be accrued. To get the 1024-QAM modulation to operate, an HFC plant with no margin will have to supply 6 dB more Signal to Noise Ratio (SNR). This begs the question: Is it easier to extend the frequency in the HFC plant by 25% or cleanup up the HFC plant for 1024-QAM? Actually, it would be synergistic to do both and gain $(1.25 * 1.25 = 1.5625)$ 56.25 % more bandwidth.

MPEG 4 Part 10 (AVC) video compression

The MPEG-4 Part 10, advanced video compression (AVC) standard provides superior compression versus MPEG-2 with the same resulting video quality. The improved suppression will allow more AVC streams per RF channel versus MPEG-2. AVC, also known as MPEG-4 Part 10, MPEG-4 AVC or ITU-T H.264, is an extension to MPEG-4. AVC offers many new compression techniques to provide equivalent quality to MPEG-2 at half or less than half the bandwidth. The most effective tool for better compression is the new mechanisms provided for advanced prediction techniques. Clearly, for digital video, both MPEG Transport Streams (TS) and IP video, the additional compression saves bandwidth.

AVC offers another bandwidth saving technique. The compression of video can be coded with a Multiple Description coding technique. Rather than code the video into a single stream, the video can be coded into multiple streams. Each stream is then sent to the decompression device. The decompression device can then use all the threads to create the resulting video. However, the decompression device may not use all of the streams to decompress the video. The resulting video will have lesser quality relative to the video using all the streams. The Multiple Description technique could be used to compress the video at different pixel resolutions. The first stream would produce video for a device with 180 lines of video (PDA, HHD, ...). The second stream could then be decompressed to generate video with 360 lines of video. The next stream could then generate a resulting stream of 720 lines of video.

Each stream then is sent downstream. A device with limited screen resolution would decode the only the first stream. A device with moderate screen resolution would use the first two streams, and a device with HDTV

screen resolution would use all of the streams. The sum of the three streams would likely be larger than the size of a single stream targeted for an HDTV application. However, the sum of the three streams is likely to be smaller than the sum of the three streams required to support video across the range of three devices in this example.

The Multiple Description feature of AVC is still under academic exploration. The technology of this technique is not available today.

Digital TV and HDTV

The large screen digital TV phenomenon is here. High Definition Television (HDTV) is the consumer buzzword attached to these monster TVs. HDTV signal transportation will load the existing HFC plants. Simultaneous, Standard Definition Television (SDTV) and HDTV transmission will burden the HFC plants.

In all of the excitement of these new bigger and better TVs, the catalyst to this technology has been overlooked. The Federal Communications Commission (FCC) has mandated by 2007 that all over-the-air broadcast TV transmissions be digital. This requires the TV manufacturers to be delivering digital ready TVs well before 2007. The one-way communications specification has already been agreed upon, and the two-way communications specification to these digital TVs is therefore implied to arrive by 2007.

The digital TV mandate provides some relief and some consternation for the MSO. An opportunity to migrate to an all digital HFC plant is clearly opened with this mandate. This converts the bandwidth inefficient analog channels into bandwidth efficient digital channels. This bandwidth relief is short lived. The digital TVs will be

HDTV capable and the consumer will be demanding more HDTV programming. An HDTV video stream ranges from 4 to 5 times larger than an SDTV video stream.

The two-way digital TVs will have the opportunity to become more than just video playing devices. The advanced digital TV may offer program guides, web page capability, or text streaming services. These are low bandwidth services. The digital TV could evolve into a device with specialized gaming electronics that allow internet gaming. As the digital TV gets more complex, the bandwidth requirements of the HFC plant will continue to rise.

Next Generation Network Architectures

To begin to deal with all of the new technologies that either are destined to be implemented in a HE or simply loading the requirements for the HE, the architecture of the HE itself must be reassessed. As noted earlier in this paper, the HE has grown into a collection of systems sharing the HFC plant. To reduce both operating expenses and the expense of turning over equipment in the HE, new architectures have been proposed.

The Unified CMTS via the modular CMTS

The beginning of this new architecture is being formed with the CableLabs Modular CMTS (M-CMTS) specification. This specification is vital in the development of future CMTSs. The intent of future CMTS architectures is to allow it to become a central entity in more than just data. The next generation CMTS must be able to control, maintain, and organize all of the streams transmitted and received over the HFC plant, not just voice and data. The unification of service control is the target of the next generation CMTS.

To efficiently create access to every CPE device, the CMTS must have access to the entire HFC plant. The traditional CMTS contained QAM modulators and upconvertors and receivers embedded into the CMTS in hardware driven ratios. For the unified service providing CMTS, embedding an entire HFC plant worth of QAMs, upconvertors, and receivers would not be cost efficient. Two systems sharing the same downstream channel must be managed. Routing all traffic through a CMTS regardless of the need for the CMTS in the downstream path is not cost efficient.

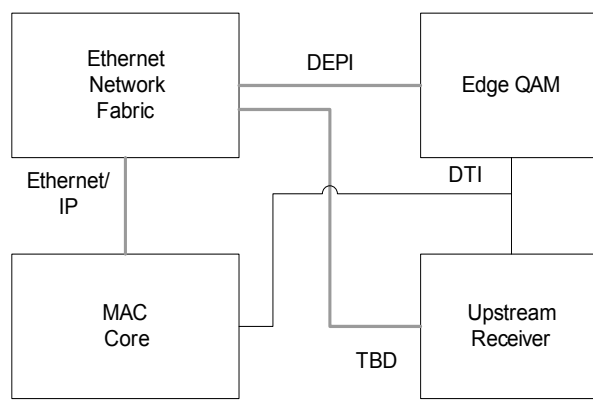


Diagram of M-CMTS

To deal with the cost and manageability of the downstream of the HFC plant the modular CMTS is required. The modular CMTS takes all of the functions of the traditional CMTS and subdivides these functions into three modules: CMTS core, Edge QAM, and upstream receiver. The CMTS Core provides all of the MAC functions required by DOCSIS. Classification, rate shaping, security, header suppression, routing, bridging, and other packet level functions. The Edge QAM performs all of the downstream PHY requirements of DOCSIS. The Edge QAM may also be accessed by other streaming devices other than a CMTS core. The upstream receiver performs the upstream PHY requirements for DOCSIS.

The Edge QAM

The key to this architecture is the Edge QAM. It is arguable whether the Edge QAM is part of the CMTS or really should be considered the stream-to-RF bridge to the HFC plant. Assuming an all digital HFC plant, the plant may have as many as 160+ QAMs on that plant. These QAMs may be transferring video, user data, voice, application information, and/or plant management information. All of this traffic does not have to flow through a traditional CMTS. However, the HFC plant may have CPE devices on any QAM that are in need of DOCSIS support. The CMTS must have access to the Edge QAMs.

The Edge QAM must be accessible by multiple devices that are providing a wide array of services over the QAM channel. The ideal Edge QAM solution would be to have an Ethernet port for receiving digital traffic and an RF connector to transmit the traffic onto to the HFC plant. An Edge QAM device may support a single QAM or multiple QAMs. The Edge QAM must be able to place the QAM signals on the appropriate frequencies, and there require some upconverting capabilities.

The Edge QAM device is really a sophisticated Ethernet Network Interface Card (NIC) for the HFC plant. This should allow the vendor community to increase the density of QAMs per RF connector, and reduce the cost per QAM. The MSO now can purchase the number of Edge QAM blades required to provide the necessary number of channels. If the plant extends the frequency range, then more Edge QAMs can be plugged in at that time.

The key to the Edge QAM versatility is the interface. The Edge QAM must provide some MPEG multiplexing capabilities. This will allow multiple devices to share an

individual QAM or QAMs. The Edge QAM must have low latency Queues for video transport. The edge QAM must also have prioritized queues for data and management message transport. The QoS capabilities provided by DOCSIS must not be hindered by the Edge QAM. The M-CMTS specification provides two separate constructs for transmission.

The first construct is an MPEG Transport (MPT) mode. This mode is expecting MPEG elementary streams wrapped in a UDP datagram. Some video processing to de-jitter the incoming video and correct the PCR is required for these streams. This processing is required by the Edge QAM. The interface also provides a transparent mode that does not get processed by the video processor or DOCSIS timestamp function. The interface also provides an MPT operation for DOCSIS packets. The DOCSIS MPT operation would have some resynchronizing circuitry to maintain timing. Behind all the MPT processing is a MPEG multiplexer to merge the streams onto the QAM channel.

The second construct is a Packet Stream Protocol (PSP) mode. This mode preserves the DOCSIS QoS features. DOCSIS MAC management message are typically prioritized ahead of user data on the downstream channel. To preserve this prioritization multiple queues that can be prioritized must exist in the Edge QAM. The PSP mode of operation also differs from the MPT mode in that the PSP flows are terminated at the Edge QAM where the MPT flows are transparent and not terminated. The combination of modes provide the Edge QAM the versatility it needs for different applications.

The CMTS Core

The CMTS Core provides all the CMTS core functions in the traditional CMTS. The CMTS core also is burdened with the channel bonding process. The M-CMTS has placed the channel bonding in the CMTS core versus

in the Edge QAM or a combination of both. The CMTS core will have an Ethernet interface to the HE network and another Ethernet interface to the Edge QAM network.

The CMTS core will bridge or route packets onto the HFC plant (via an Edge QAM). The CMTS core will classify, rate shape, police, bond, encrypt, and header suppress packets. The intent of the M-CMTS architecture is to have CMTS core implementations that are scalable. As more DOCSIS processing is required for a given HE, more CMTS cores could be added to the HE.

The Upstream Receiver

The definition of the interface between the CMTS core and the upstream receiver will be considered in the future. The downstream division was fairly simple. Leave DOCSIS in the MAC, so the Edge QAM was simple and could be used by non DOCSIS entities. Applying the same philosophy to the upstream, the interface would lend itself to be some sort of FEC block wrapped in an Ethernet packet. Then the CMTS core could rebuild the packet. However, reusable upstream receivers with non DOCSIS upstream devices may not be preferred.

The intent of modularizing the CMTS is to allow it to have access to the entire plant and become the standard control mechanism inside the HFC plant. With this in mind, it may be preferred to have the upstream receiver return Ethernet encapsulated packets with DOCSIS information wrapped around the completed DOCSIS packet. This would allow the CMTS core to be a management tool versus a packet processor.

The concept of having the Edge QAM provide DOCSIS packet processing has been dismissed for the current revision of the M-CMTS specification. However, there is merit to revisiting a different definition of the Edge

QAM. The current Edge QAM definition does not seem to take advantage of the capability of multiple Edge QAMs on a single blade. Concepts like channel bonding require the CMTS core to packet process the bonded stream and send the dispersed streams to individual Edge QAM devices. This prevents channel bonding, at the Edge QAM, from dynamically selecting the appropriate channel to send packet segments.

The Edge QAM device may actually support multiple QAMs per Ethernet port. The system would be far simpler if the CMTS core managed and routed packets, and the Edge QAM prepared them for transmission. A channel bonding mechanism in PSP mode would allow the Edge QAM to pull packets from a multi-QAM queue and put the channel bonded segments onto the bonded QAMs. This would allow the Edge QAM to optimally balance the QAMs versus the CMTS core making some crude estimation of where to send the channel bonded segments.

Putting the M-CMTS pieces together

The modularization of the M-CMTS provides architectural advantages. Another hidden advantage to this architecture is the connectivity of the parts. The CMTS Core, Edge QAMs, and Upstream receivers are all connected together with an Ethernet network fabric. Just as the Edge QAMs can be scaled and shared with other services in the HE, this network fabric can as well. By leveraging the commodity priced technology of the Ethernet network fabric, the MSO gets an inexpensive solution with excellent versatility. The technological and business success of the Ethernet network fabric clearly speaks for itself.

System Resource Manager

The M-CMTS is perfectly tuned to leverage shared Edge QAMs. In the

traditional CMTS, the CMTS managed the load on the CTMS QAM or QAMs. Now, the M-CMTS cannot really manage all of the Edge QAMs. The load on the Edge QAMs may not be known by the M-CMTS. A new device must be put into place to manage the loading of the Edge QAMs. This device will be called a System Resource Manager (SRM).

The SRM will manage the video servers, audio servers, IP server, or any server or device streaming content over the Edge QAMs. The SRM must have the capability to dynamically start and stop services, and dynamically route services to the Edge QAMs. The SRM may also need to monitor congested links in the Ethernet network fabric to prevent over subscription of that fabric. Again the SRM will have to perform the rate shaping, load balancing and manage the routing of the streams in the same manner that the Traditional CMTS provides these functions for data traffic over DOCSIS.

New Technologies and the M-CMTS

The intent of modularizing the CMTS was to create a HE architecture that would leverage all of the current and new technologies seamlessly. To demonstrate the value of this M-CMTS architecture, the aforementioned technologies will be applied to the SRM and M-CMTS architecture. As each technology is added, the SRM and M-CMTS will leverage the combination of the technologies. The result will be that the SRM and M-CMTS will have unified all of the services and technologies under one architecture.

DSG

In an ever competitive world, the cost reduction of the CPE device is inevitable. This will create CPE devices on the HFC plant with only one tuner. This tuner may be tuned to any QAM. For a HE to provide

access for DSG messages to every QAM on the HFC plant, that HE would require a traditional CMTS for every QAM. The M-CMTS has several techniques to handle this issue without a traditional CMTS per QAM.

If the DSG CPE device is a one-way device, then the M-CMTS is required to replicate all DSG messages destined for that device on all the QAMs. This requires a single CMTS core receiving and forwarding the DSG message or messages. The messages get tunneled through the CMTS core and on to the existing Edge QAMS, with the message going to each Edge QAM. The load on the CMTS core would be minimal.

If the DSG CPE device is a two-way device that allowed the SRM and M-CMTS to track the QAM the device is receiving, then the DSG server could send the message to the appropriate QAM. A point to point message sent through a single CMTS core and routed to the proper Edge QAM is highly efficient.

PVR

A simple example of leveraging the PVR with DSG over the M-CMTS would be to download program guides to CPE devices. This could be as simple as the show listings per channel for video, or music listings for audio devices.

Using the techniques listed with DSG this would be extremely simple. This would allow operational content to be stored on the PVR. This content could be updated at the MSO convenience.

Channel Bonding

Channel Bonding provide extremely high speed service to a single device. Leveraging DSG, PVR and the M-CMTS, channel bonding will be used to download a movie as an impulse purchase.

A customer orders a movie from a service provided on the CPE device. The application and the data required for the application can be downloaded and stored using DSG and PVR. Now leveraging the DSG tunnel to make the transaction, the movie could be downloaded in around 3 minutes bonding 10 QAMs. The downloaded movie would be stored on the PVR on the CPE device. The trailer for the movie could be downloaded in seconds and playing while the movie itself downloaded. The customer would have the option of watching the movie right after the trailer or viewing the movie at a later date.

The movie could be stored in the HE as an MPEG stream on a MPEG server. The DSG tunnel is used to inform the movie ordering application of the address of the appropriate servers to complete the movie ordering transaction. Once the transaction has been completed, the SRM can get the appropriate MPEG server to stream video to the appropriate Edge QAMs. The SRM can then prevent any further activity from oversubscribing any given Edge QAM. This example again takes very little processing from a DOCSIS CMTS core and simply leverages the network connectivity in the M-CMTS HE to get the video to the appropriate Edge QAM.

1024-QAM

There is no need to glorify this, just take the last example and go 10/8 times faster. Now instead of 10 channels in the prior bonding example, only 8 are required.

HDTV and AVC

The example now grows into having multiple CPE devices on the HFC plant. Some devices require MPEG-2 video, some devices can receive MPEG-2 or AVC, and some devices are AVC only devices. The

only real difference now is that the DSG tunnel is used to communicate multiple applications that in turn will make requests to the System Resource Manager (SRM). The SRM then configures the multiple video servers to stream the video to the appropriate groups of bonded channels for each service required. The number of MPEG-2 and AVC servers can dynamically change over time. The HE architecture and the CMTS tunneling commands to the CPE devices do not need to change. This leads to a cost efficient migration path from MPEG-2 to AVC in the HE. The same strategy will apply to MPEG-N in the future.

A second example of HDTV, AVC and the M-CMTS will mix HDTV and SDTV broadcast programming using AVC over the HFC. The assumption will be that the program that is being streamed is a broadcast program that can be viewed by several types of CPE devices. The Multiple Description feature of AVC can generate streams for HHD, SDTV STB and HDTV digital TVs. The SRM then directs the streams to the appropriate QAMs. If the CPE devices requiring more than 1 stream are channel bonding devices, then the streams can be dispersed across multiple QAMs. The HHD devices that require only a single stream could be instructed to find that stream on the appropriate QAM. The SDTV STB and HDTV digital TVs can be instructed to pull off the streams of interest off of the appropriate bonded QAMs. This is using bonding as a mechanism for picking up multiple streams as opposed to picking up a single stream striped across multiple channels. This system is using AVC for its compression and Multiple Description streaming benefits. The channel bonding can now be used as a load balancing tool rather than a burst performance tool. The DSG tunnel is used to direct the CPE devices to the desired streams. All of this is accomplished with the SRM and M-CMTS architecture.

IP Video

To make a long story short, integrating IP video takes nothing more than replacing MPEG video servers with IP video servers and providing the appropriate number of CMTS cores to support the IP Video through DOCSIS requirements. IP video requires more buffering, but the PVR or the never ending improvements in memory densities can both be implemented to deal with any IP video buffering issues.

The Big Ending

The big ending has already gotten to be a bit boring. As the examples continued to pile on new technologies and services, the system quite frankly easily handled them. The key to providing more services became a matter of having the source servers. The control and distribution mechanisms did not change. VoIP could be added to the mix. The boring solution is to add the required CMTS cores and upstream receivers to handle the load. If you already have idle CMTS cores or upstream receivers, then just use those. If VoIP declines and an IP gaming service picks

up, then the system replaces the voice traffic with gaming servers and gaming traffic. Again, add the server for the service that is required and the problem is solved.

Boring, scalable solutions may not make for a dramatic ending of a technology paper, but they would be a dream to the current cable operator. Dynamically changing services are an operator nightmare. New proprietary equipment means lengthy bring up time, expensive installing and maintenance costs, and lengthy Return on Investments (ROI). Installing a server for a service that is in demand today is both time and cost efficient. The ROI could be measured in weeks not years. Somehow the boring M-CMTS has become the cable operators' Holy Grail.

Wow, all of this supplied by an M-CMTS architecture. When the M-CMTS architecture is viewed from the entire HE, it becomes the unifying piece to the multiple service puzzle. This architecture allows the unification of services through the M-CMTS. Adding a SRM to an M-CMTS puts the finishing touches on the Unified CMTS, a system for all services.

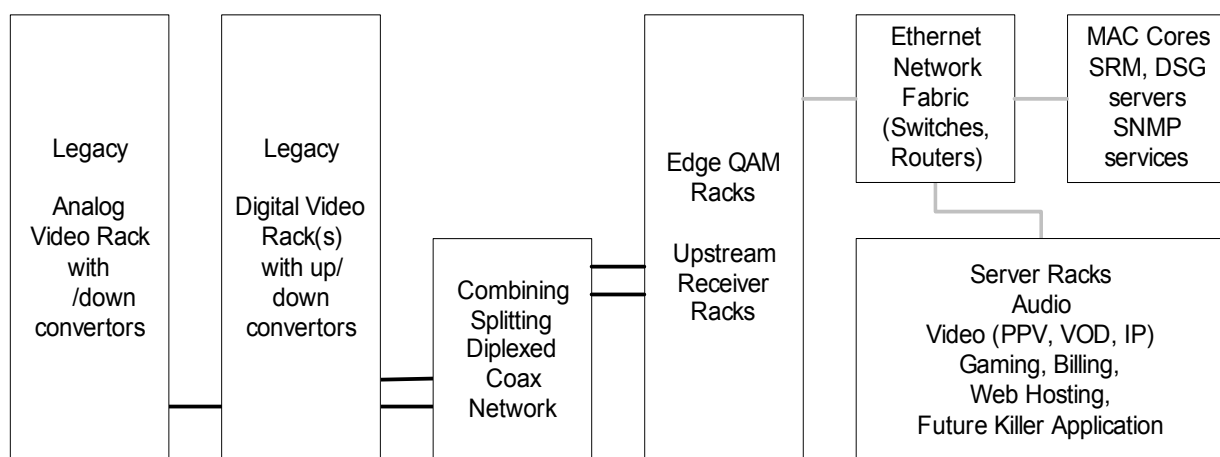


Diagram of the Unified CMTS, the system for all services

USING NEW ADI 2.0 STRUCTURES TO EVOLVE VIDEO-ON-DEMAND MANAGEMENT & SERVICES

Yasser F. Syed, Ph.D.
Cable Television Laboratories

Abstract

The new ADI 2.0 specifications allows for the separate handling and combining of the offer & title metadata, content, and its metadata. This becomes a valuable set of managing tools as catalogues get larger, new VOD marketing techniques are envisioned, and content becomes more long lasting and less traditional. This paper will first describe what are the new ADI 2.0 structures and compare this against older ADI 1.1 structures. It will also describe how VOD services are changing and what does it imply for new management and service needs and how this will be easier to handle using the new structure & backend distribution. Then it will talk about possible ways to transition from a 1.1 ADI structure to a 2.0 structure. Finally the paper will describe how the new ADI 2.0 structures can help meet these new types of demands on cable systems.

INTRODUCTION

Traditional VoD evolved as a replacement of a Pay-Per-View (PPV) model comprising of a selection of top 20 movies that are replaced every 2-3 weeks. Initial trial deployments as far back as 1993 were created to see if it could replace the PPV model. But really it was the establishment of digital cable networks that gave VoD a platform to grow on. Since then, VOD deployments among cable operators have moved beyond the trial stages towards becoming a viable digital cable service product. It is one of the differentiators

between cable services from other media service types.

Bandwidth used for Pay-per-View (PPV) services in time will be recaptured by VoD services while still retaining the value offered by PPV – current top twenty movies. But the model for VOD service can extend beyond this to carry TV shows, news, music videos, and other events found on broadcast while including possible commercial support and channel identification. This can be further extended to offer content with a continuity not offered on broadcast because of niche demands and limited bandwidth: viewing past show episodes, taking educational classes, seeing telescoping ads, touring local homes for sale, etc. The initial success of VoD bring challenges of its own that need to handle these new possibilities while scaling to increased customer usage demands.

The Cablelabs 1.1 ADI/VOD metadata documents is the current defacto specification used for delivering content with metadata to cable systems via satellite data transport or by tape delivery. Its intended purpose is to deliver a movie, its preview, and its boxcover along with metadata for a single video title offering in a one-way delivery system (see figure 1). This is the main delivery mechanism from a content provider or distributor to a cable headend or regional cable distribution center. With constant requests for new ways to offer Video content, this one logical structure has been refitted constantly to meet new VOD service, technology, and process demands; but at a considerable operational cost. This

approach becomes harder to maintain as there is more content, more ways to offer content, and shorter lead times for distributing this material.

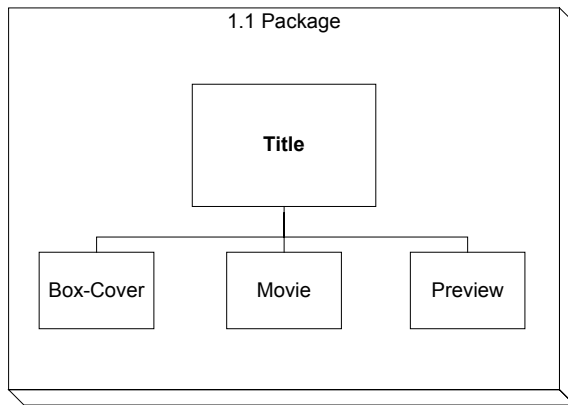


Figure 1: VOD 1.1 Structure

The new ADI 2.0 specifications create an approach to building a VoD structure more flexible to these new demands. The ADI 2.0 specifications provide a way to create more than one type of logical structure while creating a common approach to managing the delivery of these structures. The delivery mechanism has expanded to include multi-distribution approaches, more automation, and increased validation. The additional improvements in ADI 2.0 specification allow for more handling of large-sized content files and the ability to utilize this content in more than one VOD offering.

ADI 2.0 SPECIFICATIONS

The ADI 2.0 structures specification defines the basic building blocks and connectors that can be used to build logical structures for any type of VOD offering. It also describes the message document envelope needed for delivery of this information.

Assets

The basic building blocks to represent logical VOD structures are called assets. Each asset at a minimum has metadata to provide a universally unique identification (Provider ID/ Asset ID), versioning within the asset (updateNum), and management (Asset Lifetime window -- which determines how long the asset can remain with that cable operator). From an asset management perspective, the metadata is necessary to be able to track, manage, and ultimately purge assets from the cable system. There are three general types of assets: content assets, group assets, and metadata assets.

The content asset (CA) contains the content file (or link to the file) and the metadata that belongs with the content. Examples of this could be a video (movie, preview), still-image (box-cover), or even audio. This content asset by itself has no context, meaning that it is not yet associated with a particular VoD offering. This has advantages because the content asset can then be possibly shared with several VOD offerings at the same time. For example a preview could come in with a set of other previews in a barker offering and still be reused as the main preview for its movie. Furthermore it allows the option of separate transport and ingestion of large Giga Byte (GB) content files into the cable system on a time schedule different than the rest of the offer metadata. The metadata that goes along with this asset can indicate the size of the content, if it is corrupted, and whether it is completely built. Other metadata that is contained within the content asset are shareable metadata to describe the content but non-specific to the offer. An example of this would be Screen Format as opposed to BillingID.

Actions are done on the content assets through operations. This replaces the verb structure described in the 1.1 documents. “Accepting” a content asset brings the asset into the cable system as an unattached asset. “Destroying” the asset removes the asset from the cable system by effectively ending the content asset’s lifetime. “Replacing” the content asset can update the content metadata or even swap the content file though this is not often done except to recorrect existing corrupted files.

The group asset (GA) defines the context to use a collection of assets which could be a VOD movie, a movie on SVOD, a barker loop for new action movies, or whatever people possibly define. It is the organizing asset that can group and identify everything including other assets collectively specific to that context. This type of group asset has a flag to indicate itself so asset management systems can identify this as an organizing asset. This group asset can also contain its own metadata that is universally applicable for other assets it organizes.

Since the group asset is the main organizing asset indicating context, it needs to be initially ‘Opened’. Other operations for the group asset is ‘Dropped’ – the GA and its member assets are removed from system, ‘Replaced’— information in the group asset is replaced, and ‘Closed’—a hint that no more metadata or assets could be added to the GA context.

The last building block is the metadata asset (MA). It needs to belong to only one specific instance of a group asset and cannot exist on its own. It basically suborganizes a set of metadata that may be specific to only a subset of the assets organized. An example of its possible use could be releasing a title in a high definition format a week ahead of its

standard definition format release. In this case, a movie metadata asset could be created for each version containing the licensing window as well as other metadata specific for each version. The licensing window metadata cannot reside at the content asset level because it is metadata pertinent to that particular offer. An interesting aspect of suborganizing metadata into different assets is that a different provider can potentially create each asset. This can allow for 3rd party contracted contributors authorized by the provider of the owning group asset. Contrary to this, the 1.1 specifications restrict all metadata and content to originate from the same provider.

The metadata operations are in reference to the particular owner group asset. The metadata asset can be ‘Added’ to that group asset (once the group asset is ‘Opened’ of course). The metadata asset can be ‘removed’ from the group asset. Lastly the set of metadata in that asset can be ‘Replaced’ as a whole for updates to the metadata asset. The metadata asset is always subjugated to the existence of the owning group asset which means once a group asset is ‘dropped’ all its member metadata assets are removed as well.

Connectors

In the 1.1 specifications, an implicit structural definition is created that restricts the logical structure to be only one general type: namely a title with a movie, box-cover, and preview (plus add-ons). This is fine for delivering a top-twenty movie, but was a rough fit for delivering other types of content or title offers. The ADI 2.0 specifications create an explicit structural definition where the relationships between the assets under certain rules can be redefined to suit the particular offer type. (see figure 2)

The group membership connection was previously talked about in the paper when discussing metadata assets. This is the only type of membership relationship allowed. It limits the metadata asset to be a member of 1 group asset instance and prevents the MA to exist outside of that context. Reasons for this is to tightly connect all elements of the offer together such that it can be managed as a single entity. This allows for example a complete deletion of an offer by simply dropping the organizing group asset without needing to see how the changes affect other similar offers.

The second logical connection tool is the reference pointer. Because of the acceptance of the provider ID / asset ID as a universally unique identifier, assets can be tracked and be referenced by this identifier. The one-way pointer reference can be embedded in the metadata of the parent asset and can be visible to the asset manager for validation or be kept hidden until the particular application processes it. This is the main method to logically connect the shareable content asset to the offer with a reasonable amount of assurance due to the near static nature of a content asset. This same reference connector could also be used between any two assets including outside group assets though the integrity risks would need to be minimized by proper logical structure development for the specific application document. By this mechanism, it is possible to apply second level of offers, create playlists from multiple titles, or allow for other ways to innovate in the future.

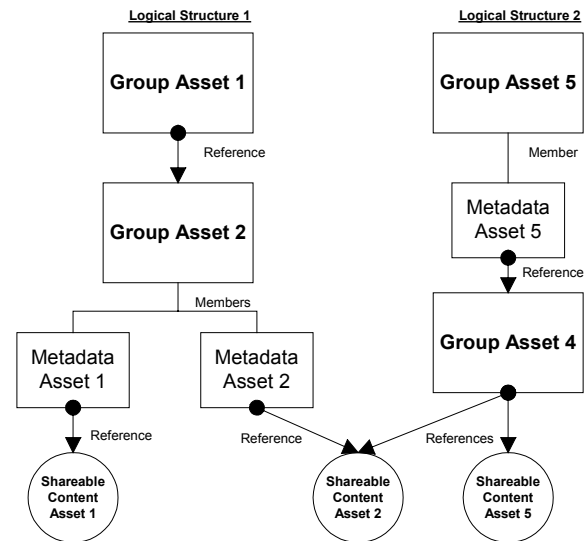


Figure 2: Explicit Structural Connectors

There also exists another connection tool (associate content signal) that is not required for creating the logical structure but does consider transportation timing issues for retrieving large sized content assets. As noted before, the content assets do not come at the same time a group asset is opened or when a metadata asset is added. Often they already exist in the system, but a copy needs to be retrieved as the offer is created. This requires preplanning to transport the file across sometimes very large regional networks. This connection tool does not create additional logical structures but ensures that the content assets are physically available to be connected with the offer.

Message Document

The message document is the delivery 'parcel' that arrives at the cable asset management system. The message document can contain operations and their assets or operations for existing assets in the cable system. Contrary to a 1.1 behavior, these operations do not have to be related to the same offer. This can allow an encoding house to pitch a batch of previews from one or even

a set of content providers to cable systems. Also contrary to a 1.1 behavior, the offer has the option of being created from one message document or multiple message documents. (see figure 3)

Metadata for the message document are used for transport types of issues. It allows for “sender” to identify 3rd parties pitching the message document. It also provides metadata for a unique identifier for the message document (docNumber), creation time, relative ingest priority, and contact info for transport link issues. An exception to the metadata being used just for the message document is an acknowledgement field. The acknowledgement field provides an http connection to report back basic level

validation on operations and assets enclosed within the document. This is an advantageous point to do this type of checking because it is the first point of ingest into a cable system. The basic level checking of the operations and assets are through message level xml parsing and validations mechanisms as well as content file integrity checks. Note this is not a basic validation of the organizing context but a basic validation of the components sent in the message documents. Validation of the offer happens at the application level where the logical structure is fully visible.

This separation of the message level and application level metadata allows for the message layer to be versioned independently from the application level. The asset manager can manage and route assets intended for applications that it need not be fully cognizant of. By allowing for this separation, the same message document layer could be used to deliver components for a VOD application or for an advertising server. Furthermore the use of versioning on namespaces also allows xml support of a mix of older/newer versions of application and delivery systems. This allows for many on-demand applications to be developed and supported using the common platform described in this specification.

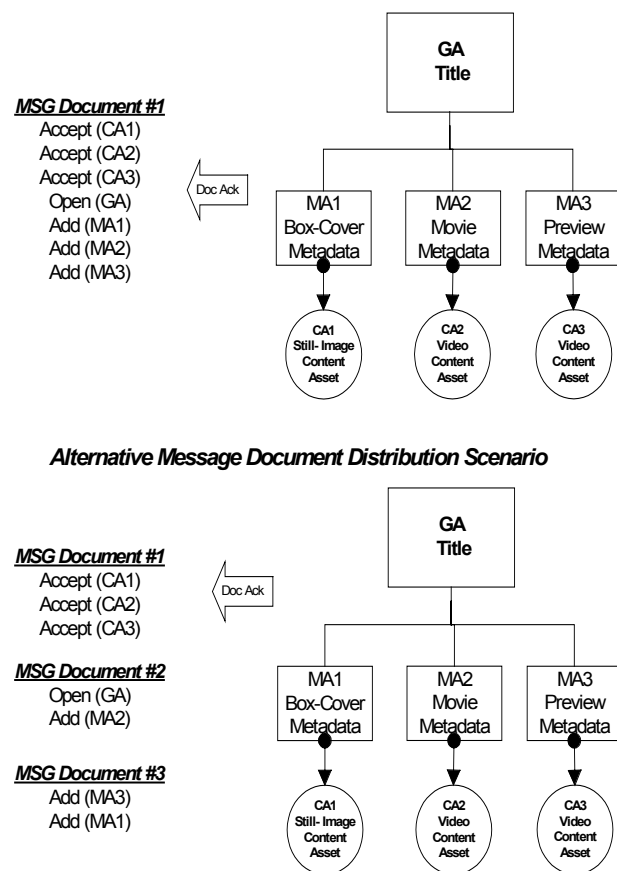


Figure 3: Single and Distributed Approaches to Create a Structure

The message document format can also be used to deliver auxiliary information to the cable systems or back to the provider. An example of this is the return ACK message (call it an ACK document) sent back upon receipt of the message document containing operations and assets (call it the ADI2 document).

This concept is extensible and has created a new specification called the ADI2.0 AIM (asset inventory messages) specification. This specification contains two types of

message documents related to VoD, but do not create any new logical structures.

The first message document described is the Notification to Deliver Content (NIDC) which is a message sent from the provider to the cable operator informing the operator on a periodic basis what content assets (movies) are being planned to be sent. The cable operator can then use this information for asset management and resource management of their servers & systems. This is especially useful as content catalogues increase and become more long lasting in cable systems.

The second message document described is the Provider Notification Message (PMN) which is basically another acknowledgement but this time on a content asset level instead of a message document level. This can be used as an event mechanism to indicate the condition of a content asset as it is further processed in the cable system.

ENABLING NEW MANAGEMENT & SERVICES

By creating a common platform to build and deliver logical structures, the ADI 2.0 specifications can assist addressing many of the concerns in transport of content and management as Video-on-Demand scales up in volume as well as new services.

Delivering content files in a large cable environment often takes time because of the transfer of large video files (in GB) across an increasingly distributed cable network. With the advent of High Definition Video, increasing the size of content catalogues, shorter delivery lead times, more demand for content storage & management, and reuse/repackaging of content; it is becoming increasingly critical to minimize the number of content file deliveries and increase the robustness of the transfer process. In this new

specification, file integrity information (filename, filesize, checksum, metadata flag), has been moved to ADI level metadata to allow non-application routers, servers, and ingestion points to check for content file integrity during transfers without need to understand the application context. A message-level ACK can indicate as soon as possible in the process if the content files need to be resent. Another content-level ACK is used to monitor the progress of file ingestion as it propagates through the cable network. For more distributed systems, the associate content signal and the content URL is useful to assemble the content components with the logical structure in time for it to be operational. In terms of shelf space management, the NIDC function gives prior knowledge of upcoming content storage demands. Lastly the content files themselves are shareable to allow them to be repurposed for other offerings and services instead of requiring the content to be resent for each new purpose.

Asset management functions benefits from the ADI 2.0 specification as well. The management of the asset lifecycle process is improved by separating out an asset lifetime window from the licensing window. This allows for situations such as having an asset exist in the system but maybe not “in-service” at that time. The specification also creates a clearer demarcation between message document functions to deliver asset components and higher-level application responsibilities. This aids in the health and fault monitoring of the delivery to headend transport network separately from logical structural issues. Also reference links can be defined that need to be recognized by the AMS to do things like validate that the referenced asset pointed to exists in the system. Alternatively, a reference link can be defined that only is processed at the application level and in the context of that

specific application processing it. Lastly, shifting to a schema-based format allows for more specific XML parsers to be used to validate elements and specific value formats of them. The acknowledgement function in the ADI 2.0 specification improves the initial 1-way pitch by indirectly providing a first level success/fail response for the pitched message document.

The new structural tools developed in the ADI 2.0 specifications can enhance at the application level the value in how assets can be offered. In 1.1 initially, the offer, title, and content were pitched and identified as a monolithic group. With the break-up of this group and new types of structural concepts introduced, it becomes easier to offer content in new ways. For instance dynamic offers can be created fairly quickly from existing content assets in the system. Examples of this are: a weekend choice of discounted films, an updated barker list for an action movie folder, a broadcasting TV show that can be offered on VoD just as it starts, coming attractions that can be updated on all VoD movies, or VoD that can be offered with or without commercials.

A new concept that can be utilized is the standardized use of a playlist which can determine the continuous play order of a list of assets. This can be used for things like sports highlights, newscasts, double features, barkers, and other things. With the creation of the metadata asset, there is an ability to have 3rd party contracted contributions as a part of the group while still retaining an overall owner for the group. For instance a music soundtrack for a movie supplied by the record company could be part of the VoD offering. Also encrypted and trick files could be added to the group asset as additional content assets that can vary depending on the VoD system that it is sent to. With these additional tools and flexibility, the variety of

possible ways to offer and organize VOD content material enables more variety of VoD types of services.

TRANSITIONAL CONSIDERATIONS FOR 1.1 STRUCTURES

The ADI 2.0 structure does not create a Video-on-Demand application. It provides the asset structure tools, and delivery transport mechanism to create an organizing asset structure for this. To create an actual VoD application based on the ADI 2.0 platform, one or more group assets, metadata assets, and content assets and their connections need to be defined. The metadata elements residing in each of these assets also need to be defined.

In creating the VOD 2.0 application for real world purposes, the existence of 1.1 structure and compatibility needs to be considered at least for a while. The VoD 2.0 structure may initially be restricted in its instance to accommodate a 1.1 transitional mechanism, but the structure itself should be flexible to allow for more this. An example of additional features could be support of a playlist to automatically update coming attractions. Another example of this is to support alternate versions of the movie within the same offer. Since 1.1 VoD systems and 2.0 VOD systems would both need to be supported at least for some time, this restricted case needs to be considered in the structural design of the 2.0 application during the transitional phase.

It is important to know that VoD 2.0 systems could still at the same time receive VoD 2.0 group asset structures that do not consider 1.1 backwards compatibility. This would allow for more creative development of services to enhance the VoD experience. It would also give incentive for older systems to transition faster from a 1.1 environment.

An example of this could delivery of a collection of TV episodes that can be offered for one price for a set or an individual price for a la carte.

Some of the immediate new features demanded in VoD services could still be deployed in this transitional phase. With the mechanism of reference pointers, new logical structures can refer to other logical structures. Since this reference mechanism points to the unique identifier of the asset, this same reference could point to a 1.1 VOD title as well. This would allow for some new 2.0 features to incorporate existing 1.1 titles. For instance, dynamic category folders like an action folder can refer to a set of existing titles including 1.1 titles. Some of these titles can also belong in other folders simply through the reference mechanism. These types of strategies can ease the transitional pains of switching cable systems to the new ADI 2.0 platform.

CONCLUSION

This paper describes the ADI2.0 specifications. It defines general asset types and their relationships to create many different types of logical structures as opposed to a single logical structure defined for VoD 1.1. It also describes how a structure can be created through one or more message document deliveries into the cable system. These new mechanisms create flexibility to address some of the existing issues with delivery of large content files, management of VOD titles, and creations of new types of VOD offerings. It also addresses some of the

scalability issues from a structural perspective as VOD increases in both volume and variety. Many different applications can use these same ADI2.0 specifications to create logical structures for there specific type of VoD offering. The most immediate application to be developed on the ADI2.0 platform will address the single title VoD offering that allows for a transitional phase with existing 1.1 VoD platforms.

REFERENCES

1. Cablelabs Asset Distribution Interface Specification Version 1.1, www.cablelabs.com
2. Cablelabs VoD Content Specification Version 1.1, www.cablelabs.com
3. Cablelabs ADI2.0 Structures Specification, www.cablelabs.com
4. Cablelabs ADI2.0 Asset Inventory Messages Specification, www.cablelabs.com
5. "VoD Victories: Be Careful What You Wish For", Carl Weinschenk, Comm. Tecnology, Feb. 2004.
6. "Supporting Large Scale VoD Through Common Metadata", Yasser F. Syed, Proceedings of NCTA 2004, New Orleans, May 2004.
7. "Game Planning For Bigger, Better VoD", Karen Brown, CED, May 2004.
8. "Ingest & Metadata Partitioning: Requirements For Television on Demand", Robert G. Scheffler, Proceeding of NCTA 2003, Chicago, May 2003.
9. "Video-on-Demand:Quality Shows", The Hollywood Reporter, December 1st, 2003.

VOD SERVERS – EQUATIONS AND SOLUTIONS

Glen Hardin¹ and W. Paul Sherer²
¹Time Warner Cable, ²Arroyo Video Solutions

Abstract

Video-On-Demand (VOD) is now a widely deployed product with a ready audience. No longer a “trial” product, it is a cornerstone offering for the cable industry - generating revenues, reducing churn and setting MSOs solution apart from satellite.

Yet the technology underpinning VOD services is still in its infancy, and, as new VOD services are developed, the VOD infrastructure must continue to evolve if the potential of these new services is to be realized to its fullest.

This paper seeks to provide context for VOD server technology - where it has been, where it is and where it might be going. This discussion is presented in context of the changing VOD server equation. Understanding this equation is paramount to understanding the solution going forward.

THE ORIGINAL EQUATION: CONTENT + STREAM = VOD

Historically there have been two basic variables to the VOD equation: the content variable and the streaming variable. Each VOD server solution has attempted to understand and resolve the relationship between content and streaming. All vendors in the marketplace work to optimize performance and price as they tackle the basic problem of how to access the stored content, transfer it across the bus architecture and pump it out of the video server without interruptions. Some do it with brute force and

others with complex elegance. At first glance, it is seemingly a basic problem to solve, but, as the multipliers in front of each of these variables scale independently and infinitely, the solution quickly becomes complex. The physical solution to the equation can be based in either proprietary or commodity hardware and bonded together with plenty of custom software

Content:

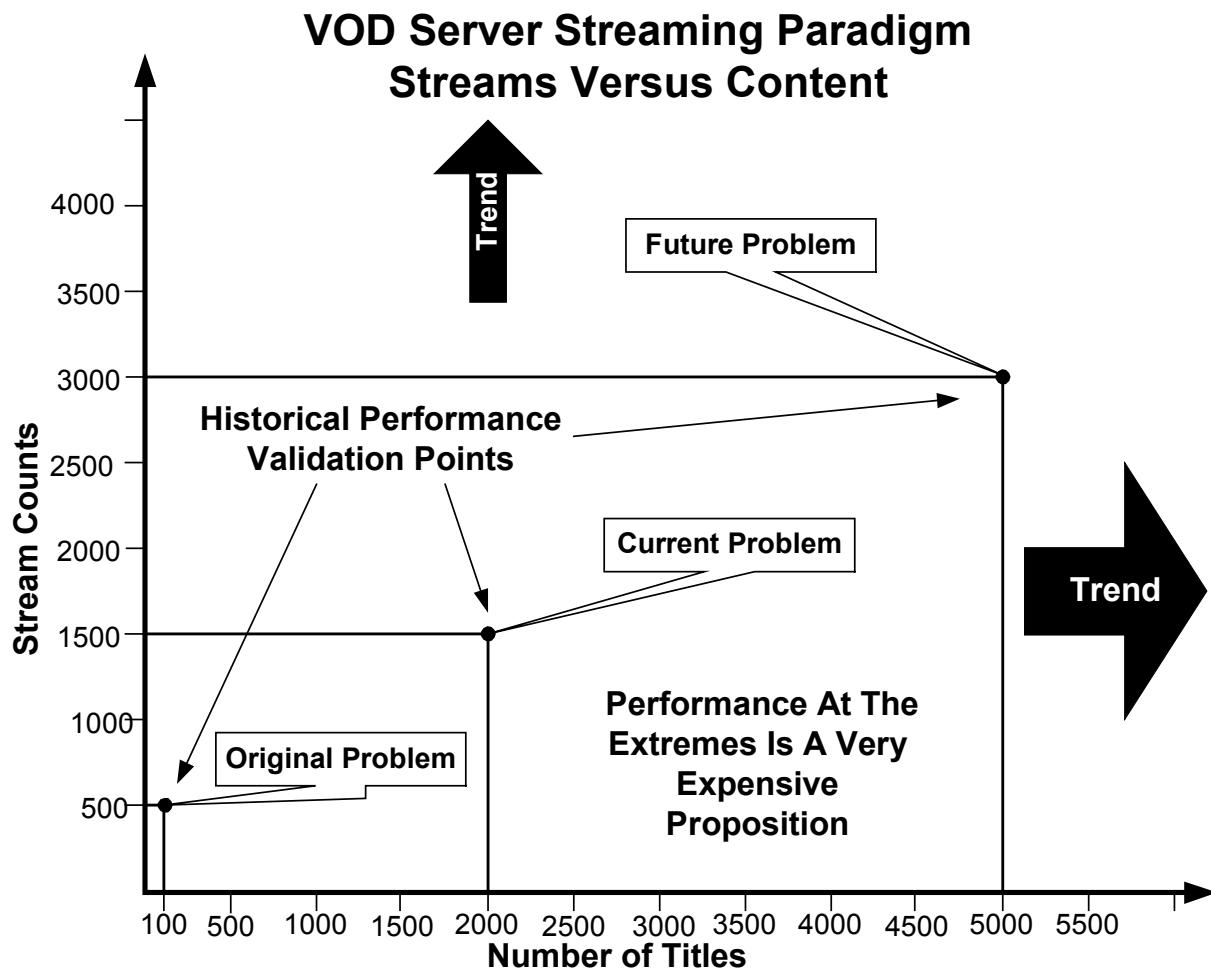
The old real-estate adage “location, location, location” has its analogy in VOD and it is “content, content, content”. Content is the main driver for the success of VOD. Add additional compelling content and the stream use rates will increase.

The amount (and type – High Def content is 4 times as resource consuming as Standard Def) of content drives the total amount of storage the system requires. In the original VOD services, the content variable was limited to the top 100 “hit” titles – requiring perhaps 250 content hours of storage. As VOD technologies proved themselves, new services such as subscription video-on-demand were added and the total number of storage hours grew to support them. The hours of storage grew from a few hundred hours to 800 hours. With the increase in the number of subscription services and recent new services such as Free-On-Demand, Music-On-Demand and High-Definition-On-Demand (HDVOD), the storage requirement quickly has quickly grown to thousands of hours.

Depending on a server's architecture, scaling the content storage may be as easy as adding more drives or an additional disk array to the system or as complicated as replacing all the drives in the system. The one thing that is for sure, if the VOD service is to be successful, the multiplier to content variable can go in only one direction, ever increasing.

everything and it must be accomplished flawlessly without interruptions.

Scaling streaming is a very complicated proposition and different vendors have approached the problem in different ways. Historically, VOD server vendors relied on core disk Input/Output (I/O) subsystem performance to attain their stream



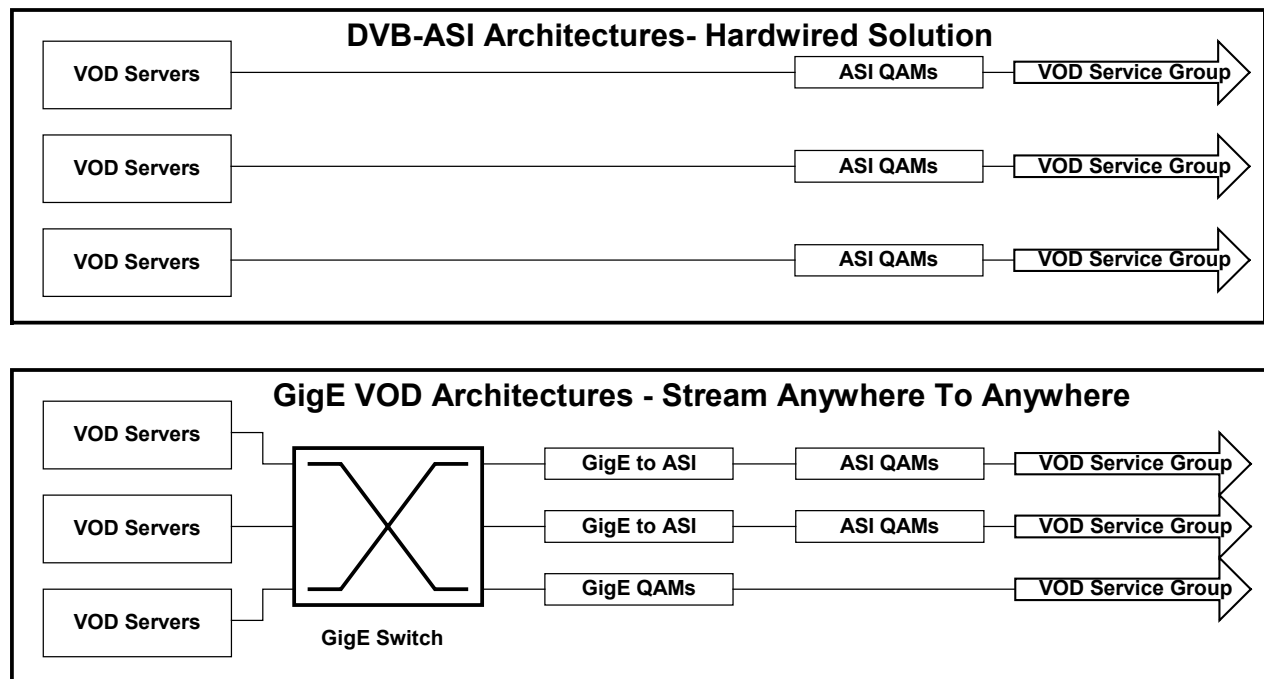
Streaming:

Streaming needs to access the content stored on disk and route it across a bus or interconnect and pump it out of the server. Since the sole purpose of streaming is to deliver content, all other functions may need to operate at a lower priority including the reception of new content. Content delivery is

performance. Some vendors chose to implement complex interconnect and RAID architectures to gain the efficiencies of parallelism and thereby increase streaming performance, while others scaled through simple server replication. In either case, validating a server's streaming performance was accomplished by taking a single piece of content and streaming it out at the server's

max stream capacity (the easy way) and taking a unique piece of content per unique stream to the server's max capacity (a much harder problem to solve). Testing at both extremes guaranteed that the server could deliver the content in any way the customer could ever order it. This performance at both extremes came at a relatively high price, but

into the VOD server platform, and immediately (within seconds) allow all customers to stream them. This rapid increase in ingest requirements is a natural outgrowth of the increase in content offered in VOD form, but it also seems to be a universal in the various next generation On-Demand services under development - including



was reasonable with a relatively small amount of content.

THE NEW EQUATION: INGEST + CONTENT + STREAM = ON DEMAND SERVICES

As discussed above, server architectures have historically focused on optimizing the output capabilities of their servers at the expense of their input capabilities. However, increasingly a new factor is changing the original server performance equation. The new factor is ingest.

Ingest:

Servers are increasingly required to receive MPEG files in real-time, ingest them

Network Personal Video Recording (NPVR), broadcast "Start-Over" and client applications like Weather-On-Demand.

These real-time acquisition-based services greatly impact VOD servers and in multiple ways. Content storage requirements are growing tremendously as the number of networks offering On Demand content grows. Instead of supporting 1200 titles, the VOD servers increasingly need to support multiples of that number. Streaming is also impacted, both because the wealth of new content must be written to non-volatile storage (i.e. disk), and because of the increase in the quantity of streams as subscribers access the new content. Additionally, since the quality of service must be maintained both for content ingest

and for streaming, VOD servers will have to work within even tighter performance tolerances as both these variables scale. This equation is far more complicated equation than what was originally required in the early days of VOD.

Architecturally, real-time acquisition-based services favor more centralized content storage solutions that allow single ingest points to serve all customers. The ingest server must have interconnectivity to all service groups. Supporting such features in highly distributed server architectures is overly complicated and almost impossible.

With the broader width of content offering and the advances in parallel technologies, the current VOD server architecture paradigm needs to be re-examined. Furthermore, the market now has the historical experience to evaluate the necessary performance requirements against the usage patterns of the On-Demand-Services offered.

REMOVING THE COMPLEXITY FROM THE VIDEO SERVER

Advancement in other technologies, including software technologies, has allowed the complexity of the VOD server to be simplified.

ASI to GigE

The most important shift in complexity of the VOD server was the removal of the DVB-ASI interface and replacement with the GigE interface. As a result, the VOD vendor no longer had to develop and support custom DVB-ASI cards within the server, which was a huge cost reduction for the server companies. VOD servers with DVB-ASI also require the video server's streaming capacity to be in parity with the edge capacity as the video servers are physically tied to the edge

devices. This shift to GigE also reduced the barriers to entry, allowing new vendors and innovation into this market.

The Edge

As a result of the shift to GigE within the video server, the requirement and costs for DVB-ASI moved further out into the network. To maintain compatibility with the existing QAM devices, new devices were developed to translate the GigE back to ASI to interface to the existing QAMs. Now, native GigE interfaces are available from every QAM manufacturer, negating the requirement to convert the GigE signals back into ASI prior to the QAM. This will further reduce the cost and complexity of the VOD solution.

Transport

Advancements within the transport technologies have greatly facilitated the shift from highly distributed VOD architectures to more centralized architectures. Transport technologies have gone from inefficient ASI transports to single GigE pipe on a pair of fibers to 40 times 1G, 40 times 2.5G and finally 40 times 10G on a single pair of fibers. Highly distributed architectures also required multiple instances of storage arrays and copies of the content. Centralizing storage and/or the servers has the added benefit of allowing for greater efficiencies through sharing of the storage arrays across many streaming devices. As a result, fewer storage arrays are required as fewer copies of the content are needed.

Core Switch

The most desirable and advantageous method of connecting the GigE video server into the cable plant is through a GigE core switch. Advancements in switching

technologies now allow for fully meshed non-blocking delivery of video.

The integration of a core switch enables video servers to stream from anywhere to anywhere. That is to say, any video server streaming port can service every VOD service group. Since content is no longer tied directly to storage at the edge, but is now centrally available to all streaming devices, the content is no longer bound to a particular video server component. The core switch also has the added benefit of reducing the overall streaming requirement of the server. The server's streaming capacity no longer needs to be in parity with the edge QAM capacity, but only with the max stream utilization. This allows the video server to scale independently from the edge QAM capacity. For example, a DVB-ASI server with the capacity of 3000 streams serving a given

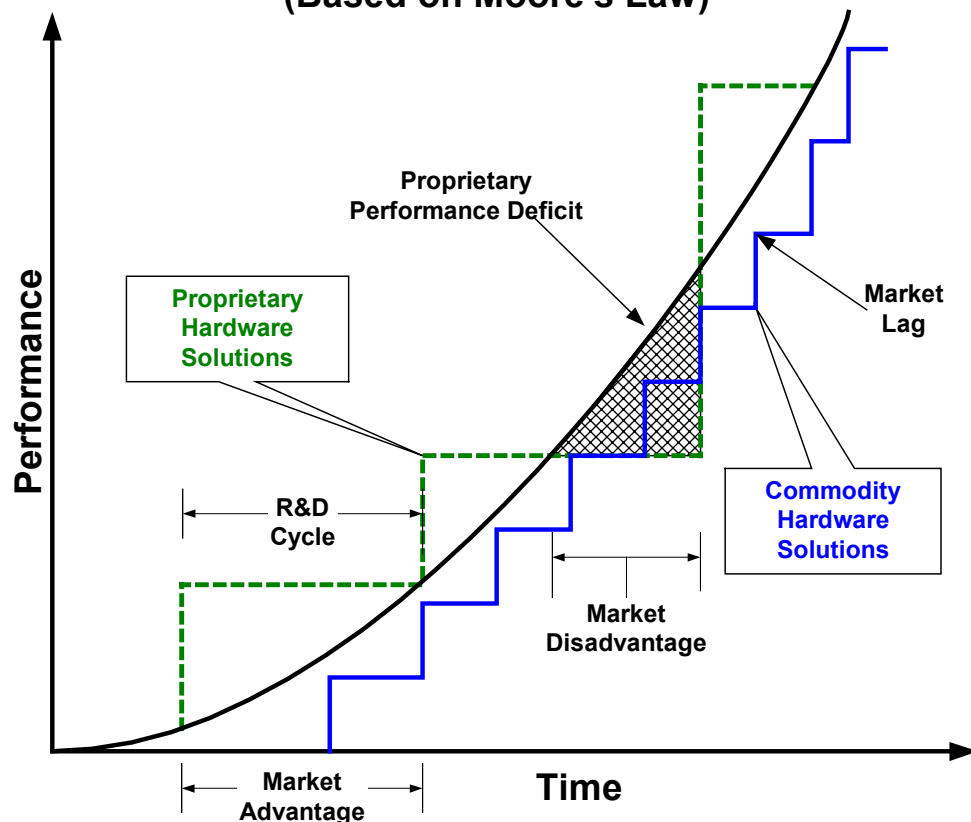
customer base can now be served by a 2000 stream GigE server with the same blocking factor.

Furthermore, the addition of the core switch between the VOD server platform and the plant also minimize the requirement to develop methods of interconnecting various discrete storage arrays though some backend back-end switching fabric.

Software Infrastructure

Finally, advancements in open software standards that allow interoperability between VOD vendors have greatly influenced the marketplace. MSO's are no longer held captive to a particular vendor once the initial purchase is made. Each time the system expands or major features are added and new server capacity is required the MSO can

Commodity Hardware Performance Curve (Based on Moore's Law)



choose the best of breed among the vendors. This ensures that VOD vendors remain competitive in terms of price and performance.

PROPRIETARY HARDWARE SOLUTIONS VS. COMMODITY

The graphic attempts to illustrate the relationship between video server performance and commodity hardware performance based on Moore's Law, an industry-accepted concept that hardware performance will double every 18 to 24 months. In general, the hardware commodity performance curve increases due to parallel advancements in all the technologies within the PC market: faster and multiple processor machines and bus infrastructures, faster and denser DRAM, drive technology, and network interfaces.

The increased performance described above results in several secondary benefits for the constrained environment of the cable headend: less space, power, cooling, and wiring are required. The newer solutions are much more dense and efficient in terms of Mbps per rack unit and the number of Mbps per watt of power consumed. Less power consumed infers lower cooling requirements. Additionally, as the outputs of the video server become denser, fewer wires are required to integrate these servers into the plant. It has been demonstrated that a server of 5000 streams @ 3.75 Mbps can now be wired into the plant with just a couple of 10 G interconnects. Historically, this interconnect would require upwards of 31 DVB-ASI wires or even 21 Gige connections. Less wiring substantially simplifies the integration work.

Both proprietary and commodity server vendors try to optimize their server costs because the stream price is determined by competition within the market. A vendor can

only control its server costs. If, for a given set of hardware, a video server can sustain n number of streams, then the minimum per stream cost equals $\$ / \# \text{ streams per unit of hardware}$ (not including the cost of development for the necessary software and other associated costs). Regardless of the hardware solution, the software is the valuable component of intellectual property of any vendor.

It is prophesized that this curve cannot sustain its exponential growth forever but in the near-term it provides guidance and insight into the future capabilities of the market.

Proprietary Hardware Solutions:

Many VOD server vendors have developed a proprietary solution by creating and integrating custom hardware components and/or custom interconnection technology. If accomplished effectively, the resultant solution should outperform what is available in the commodity market using a similar generation of technology.

The difference between the performance of proprietary solution and the commodity curve determines the performance advantage of the proprietary solution. The performance advantage translates into a market advantage for a period of time until the commodity curve catches up with the proprietary performance. Server companies offering proprietary solutions must exploit this finite time of market advantage through sales to recoup their investment in hardware R&D. At the same time they must also continue to invest in the next generation server solution lest they fall below the commodity performance curve. It is a never-ending race to stay ahead of the commodity curve and a risky business proposition. It is easy for vendors with proprietary solutions to fall

below the commodity performance curve if they do not carefully time their adoption of the newer higher performing hardware. There is a high cost to develop performance gains above the commodity performance curve leading to expensive R&D cycles. Those R&D costs must be re-cooped before commodity performance catches up, otherwise, sales opportunities will evaporate as the performance advantage disappears. In sum, it is possible to develop a proprietary solution that exceeds the commodity curve, and the more is invested the longer this advantage will remain. However, the commodity market has proven time and time again that the commodity curve will eventually catch up regardless of the technology.

Since MSOs cannot be expected to perform forklift upgrades enthusiastically or frequently, proprietary hardware solutions also face the challenge of integrating newer higher performing hardware into an existing lower performing solution. Typically, proprietary solutions rely on symmetric server performance with all machines within a server complex operating at the same performance level. But it does not make sense to integrate new high performing hardware and operate it only at the existing performance levels. Therefore in addition to constant efforts to keep up with the commodity hardware curve, vendors of proprietary solutions must undertake the development of many lines of custom code in order to have older and newer generation hardware interoperate (if at all possible) at their respective levels of performance as one integrated seamless solution.

Hence, even in proprietary hardware solutions the software is as important as the hardware and is, in fact, the key intellectual property within the VOD server platform.

Commodity Hardware Solutions:

In the emerging VOD industry back in the early 90's, the raw commodity server market barely was able to eke out enough performance from a given platform to justify the costs of VOD. The market price for streams was magnitudes higher than it is today.

There are several enabling factors that allow for commodity hardware solutions to be competitive today. First the content equation has changed dramatically as described above, i.e., 10,000 hours of content versus the historical 100 hours of content. Second, the base hardware available in the commodity market has the necessary off-the-shelf performance required to deliver dense VOD streaming. Video servers supporting multiple GigE and 10 GigE pipes per two or three rack units. And finally, the MSO market has accepted the premise of caching based on its historical content use patterns and the cost/performance trade-offs associated with cache-based servers.

VOD vendors with architectures based on commodity off-the-shelf servers abstract the hardware solution from the software solution and, at a minimum, develop loosely coupled systems. The VOD delivery solution is software-based, which makes the hardware choice an independent decision. As such, the vendor is able to choose the best-of-breed within the commodity market.

A potential drawback to working solely with commodity hardware is that the performance of commodity platforms must lag slightly the commodity performance curve, due to the need to re-qualify new hardware platforms as they become available in the market. Best-in-class solutions abstract out the software from the hardware allowing

performance to more closely follow the commodity performance curve.

Vendors using commodity servers must, like those with proprietary architectures, address the problem of integrating advancements in hardware into their solution. To address this issue, a commodity based solution needs to be developed in a manner that supports asymmetric server performance within the solution.

By achieving solution independence from the underlying hardware, commodity vendors allow MSO's to utilize existing procurement and maintenance contracts for the underlying hardware. This allows the MSO to leverage its volume purchase agreements with commodity hardware vendors. Further, internal expertise can be more readily leveraged across hardware platforms which are used for multiple service solutions.

Another reason to focus on commodity hardware is enhancing the utility of the intellectual property which must be created by the vendor. Simply put software intellectual property is readily reapplied across multiple generations of underlying commodity hardware with little to no redevelopment or retraining across generations. This allows the commodity vendor to focus on continual enhancement of system robustness and functionality without the need for continual investment in long lead time hardware development cycles in order to keep pace with the performance created by the overall computing market.

All servers, even the historical servers, stream from RAM. The difference with the new caching servers is that the RAM is used to capture the "working set" of the cache. That is, the set of content which is active at this given moment and likely to be active in the near future. The goal is to minimize the

size of the working set so the minimal amount of expensive components may be utilized to achieve the desired level of performance.

CACHING SERVERS

Overall caching server vendors try to optimize the right mix of components and costs to get the greatest return on performance

Caching components vs. Costs

RAM

~\$350 / GB

High Performance Drives

~\$5 / GB

Standard Performance Drives

~\$1 / GB

When a caching architecture is part of any system be it a microprocessor or a VOD server one of the first questions which must be asked is how to determine what to keep in the cache and for how long. All such systems use a mixture of predictive and reactive algorithms to decide what to cache.

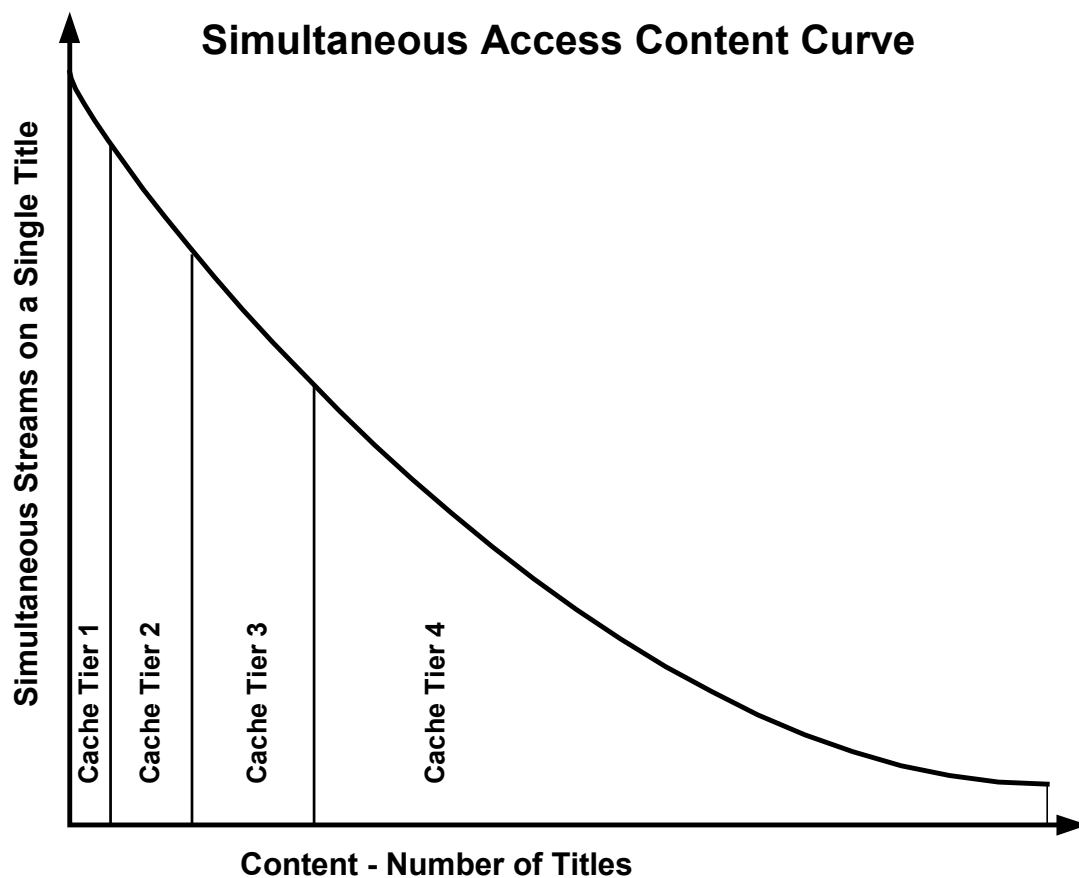
The most common predictive algorithm is the "next obvious thing". That is based on whatever is happening now the next obvious thing will most likely happen next. In a microprocessor, this usually means the next instruction after the current one – in a VOD server this usually means the frame after this one. Some seek to predict events at a much higher level. In a microprocessor this might be to predict which program will be run or in a VOD server which title will be played. The problem with this approach is the decision at this level depends on factors which are beyond reasonable prediction a priori. Johnny

Carson can die and the "Best of Carson" titles can suddenly become very popular. Janet Jackson can suffer a wardrobe

malfunction and a particular sequence from the Super Bowl can see extremely high utilization.

Most effective caching systems do not rely on accurate prediction at this level but rather rely on reactive algorithms. The approach is to make observations at a slightly broader level than the low level predictions described earlier. But to then assume that these higher level decisions will be tend to be self-similar. That is, for example, if 50 of the last 100 plays have been for a certain sequence of a Super Bowl then is it is likely that many of the upcoming plays will be for the same sequence. In this way a cache can adjust quickly to decisions which it cannot predict accurately beforehand but it can observe accurately and react to.

Traditional VOD servers tended to treat a piece of content as whole but the “wardrobe malfunction” example illustrates a portion of a piece of content may have radically different usage patterns associated with it than other portions. A more efficient cache can recognize that this portion has radically different usage and treat it differently than the rest of that piece of content. These usage patterns can happen for many reasons in addition to an event in the content. New forms of navigation such as chaptering can allow entry into a piece on content at a set of locations. The chapters in effect become mini-pieces of content with a larger whole. Another trend is the creation of virtual assets. These are logical assets composed of components from other assets which are perceived as single asset by someone viewing



them. This could be an edited topical news update or a play list of music videos.

Another reason to no longer view a piece of content as whole is responsiveness. Even if a user has held a bookmark for resume for a long enough time that the local cache has flushed the content, it is desirable that the VOD system should be able to “resume” very quickly and easily. To do this, the caching system should retrieve content starting from the point where the resume occurred, rather than from the opening scenes of the piece of content.

As has been noted earlier, VOD is now a cornerstone revenue generating service. As such it must be robust and available 24 X 7. This means MSO’s should look to vendors to provide automatic resiliency to system faults and to allow for maintenance and upgrade without service outage.

One important consideration is the unit of failure for which the system is resilient. When considering the failure modes which must be compensated for, most would think some hardware fault such a network interface failure. In reality, for all types of video server, whether based on proprietary hardware or commodity hardware, the most common failure mode is a failure in the software not a failure in the hardware. So in this sense the most common unit of failure must be considered to be the server itself. This means the entire server function must be recoverable automatically. That is, the current workload must be recovered intact by other systems without the need for human intervention. This level of resilience has been applied to telephony and data applications but is just now being designed into VOD servers.

In terms of content availability, while there is a definite trend to centralized storage,

many MSO’s are now considering geographic resiliency in their system planning. Centralized but in at least two locations and interconnected through switching and transport. The idea is that the content storage must survive a natural disaster such as a hurricane or tornado or a manmade problem such as a portion of the power grid going offline. By having content stored in geographically diverse location the odds that such an event would take two or more facilities offline is greatly reduced compared the odds of a single facility going offline.

For many years the resiliency of content was assured via RAID 5 technology. With costs of modern disks becoming so low, in many situations it is simply more cost effective to keep multiple copies. The issue with RAID 5 resiliency structures is that there is an assumed extremely high bandwidth path among the components of a RAID 5 structure. This is reasonably easy to achieve among components in single system but becomes increasingly onerous when resiliency is spread across many systems. In a RAID 5 system of n components when a failure occurs, all n-1 remaining components must participate in recovering the lost information - which must be regenerated through computation. The bandwidth impact of this process will often make it impractical particularly across geographically diverse content storage facilities.

All of the above is leading to the creation of caching tiers – each with a role to fulfill. The exact boundaries of these tiers will to large degree be determined by the cost and reliability of transport between the tiers.

Level 1 – The Edge

This is the tier closest to the service groups. The content kept here will be fairly

active and will have fairly high reuse. The two more expensive caching components of RAM and high performance disk drives will be used in this tier. The key here is to capture the “working set” with the minimum amount storage capacity. The goal of this tier will normally be to satisfy 90 to 95% of the stream requests within this tier to provide sub-second responsiveness. However, while handling the bulk of the stream, this tier would have very little of the overall cache storage capacity - typically only a few percent. This creates a very efficient usage of the high cost caching components in this tier.

Level 2 – Local Storage

The tier behind the edge serves to decouple the higher instantaneous performance of the edge from the much more modest performance of the local library. This tier can be seen as a performance matching tier which uses greater cache storage capacity to only allow a small number of the total stream requests to have impact on the local library. At this level the storage is still viewed primarily as cache with the implicit assumption that if need be content can always be retrieved from the local library and resiliency of content is less important.

Level 3 – Local Library

This tier is the demarcation of the relatively inexpensive and readily available local transport to the relatively expensive and scarce long haul transport. This content has high resiliency so that single failures of devices or servers can be handled without requiring retrieval from the regional or national library. If the area served by the local library is large or prone to problems such as hurricanes the storage may be implemented with geographic resiliency. The percentage of all content accessible from the associated edge systems is very high.

Because of the large amount of storage lower cost storage components are used here. It would be expected for every stream play request which accessed the regional or national library that thousands or tens of thousands of stream play requests would have been seen by the edge systems.

Level 4 – Regional or National Library

This tier is the ultimate source of content available to any edge system. All content available to any edge system is resiliently stored somewhere in the library. This tier will have geographic resiliency and multiple points of ingest. The regional or national library tier has the greatest storage capacity and the greatest resiliency of all the tiers.

NEW TESTING PARADIGMS

The advent and adoption of new caching server architectures requires a re-examining of how servers are tested and qualified. The old method of validating a server’s streaming performance by taking single piece of content and streaming it out at the server’s max stream capacity and by taking a unique piece of content per unique stream up to the server’s max will not result in the desired and practical price and performance point. The historical usage patterns must be applied to the testing and validation of the new caching servers.

A new term needs to be defined to help normalize the validation of servers. The term Cache Gain represents the additional steaming capacity above what is available through the core disk I/O subsystem performance. For example, if a given server has a disk I/O performance of 1000 unique streams but can deliver not only the 1000 unique content streams from disk but an additional 500 duplicate content streams from cache, the server would demonstrate a Cache Gain of 50 percent. In that same example, if

a server were able to deliver 1000 duplicate content streams from cache then the Cache Gain would be 100 percent.

Since caching servers vary greatly in the number of cache tiers and performance within the tiers, setting simple easy standards of performance is difficult. The process of normalizing the system performance to core disk I/O performance provides a baseline from which to work.

CONCLUSION

The jury is still out. Although all the best and brightest within the VOD server community agree that there are cache gains to be made as of now, there is not enough empirical data of cache effectiveness to unequivocally say what exactly what the cache gains are for a given type of content and service.

What is a believable cache gain? Is it 10%, 100% or 1000%? Only careful monitoring of live systems in the field will prove out the actual achievable gains.

However, this architectural advance clearly represents the next step in the evolution of the infrastructure for on demand services. With the advent of this architecture the stage is set for a plethora of new services reliant on much greater breadth of content and much more dynamic usage. "Start Over" and network PVR fit well to this architecture. More applications will come.

One can now see the infrastructure coming into being which will enable the efficient delivery of a fully personalized entertainment to every MSO customer on every television.

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.