

The Internet Protocol Journal

June 2012

Volume 15, Number 2

A Quarterly Technical Publication for
Internet and Intranet Professionals

FROM THE EDITOR

In This Issue

From the Editor	1
Transition Space	2
December in Dubai.....	17
IP Fast Reroute	30
Letters to the Editor.....	35
Call for Papers.....	39

Deployment of IPv6 took another step forward on June 6, 2012, when numerous website operators, *Internet Service Providers* (ISPs), and home router vendors participated in the *World IPv6 Launch*. Organized by the Internet Society, the event attracted significant media attention as the participants enabled IPv6 permanently and rendered it “on by default.” More information about the event is available from www.worldipv6launch.org

Migration to IPv6 is not a simple task, as outlined in many previous editions of this journal. Various tools and techniques have been developed, one being the use of so-called *Carrier-Grade NATs* whereby the end customers connect to the Internet using private (RFC 1918) addresses and the ISP provides translation for both public IPv4 and IPv6 addresses. In April of this year, the *Internet Engineering Task Force* (IETF) approved and the *Internet Assigned Numbers Authority* (IANA) allocated a new IPv4 address block (**100.64.0.0/10**), designated for use as “Shared Transition Space” in support of the IPv6 transition. We asked Wesley George to describe the rationale behind the use of this additional private address space and discuss the debate that resulted from this allocation.

The world of telecommunications has changed dramatically as a result of the rapid expansion of the Internet. Traditional telephone lines are being replaced by *Voice over IP* (VoIP) systems for both private and business use. These changes represent big challenges for traditional telephone carriers, and even for some countries whose income used to depend largely on telephone “settlement charges” for international phone calls. The *World Conference on International Telecommunications* (WCIT) will take place this coming December in Dubai. Geoff Huston discusses some of the proposed changes to the *International Telecommunication Regulations* that could affect the Internet in various ways and will be discussed at WCIT.

The IETF is concerned not only with IPv4-to-IPv6 migration, but also with recovery upon router or link failure. In our final article, Russ White describes *IP Fast Reroute*, a technique for providing fast traffic recovery when these failures occur.

As always, your feedback about anything you read in this journal is most appreciated. Please contact us at ipj@cisco.com and don't forget to renew your subscription and provide us with any postal or e-mail changes.

—Ole J. Jacobsen, Editor and Publisher
ole@cisco.com

You can download IPJ
back issues and find
subscription information at:
www.cisco.com/ipj

ISSN 1944-1134

Shared Transition Space: Is it necessary?

by Wesley George, Time Warner Cable

Recently, the *Internet Engineering Task Force* (IETF) approved^[1] and the *Internet Assigned Numbers Authority* (IANA) allocated^[2] a new IPv4 address block (100.64.0.0/10) designated for use as “Shared Transition Space” in support of the IPv6 transition. This decision was highly controversial within the different standards and policy bodies that discussed the idea. The author would like to note that people have been debating this topic for years, and nearly everyone within the broad stakeholder community seems to have a strong opinion on the matter, including me. Despite the best of intentions, some of my opinions and biases may appear within the article. I did not intend this article to be a definitive conclusion on the matter, but rather a summary of the recent discussion. Whether the standards bodies involved came to the “right” or “wrong” conclusion—as well as the veracity of the arguments on both sides—is an exercise for you, the reader.

Internet Service Providers (ISPs) and users have significant investments in equipment and applications that must be updated to support IPv6. Progress is accelerating with regard to IPv6 availability in hardware, software, and access, though broad availability remains a long-term problem. In the interim, IPv4 will continue to be an important capability for providing users with access to Internet resources. As a consequence, considerable effort has been expended in conserving the increasingly scarce IPv4 resources while maintaining “business as usual.” This conservation has taken the form of policies for address allocation and management^[3], as well as new protocols and technologies. It is likewise important to note that ISPs must manage IPv4 exhaustion in a way that is least disruptive to users while undertaking full IPv6 deployment—two completely different and parallel activities. Any business that relies entirely on efforts to extend the useful life of IPv4 without executing on an IPv6 deployment plan is merely delaying the inevitable effects on their customers and ultimately their profitability.

IPv4 “life extension” is an area that remains controversial. Some believe that any effort to extend the useful life of IPv4 and allow the IPv4 Internet to keep growing beyond its original design limitations will seriously affect the timeliness of reaching critical mass with IPv6. The idea that many opponents of the “life-extensions” methods are supporting is that IPv4 exhaustion and the resulting transition from IPv4 to IPv6 is going to be disruptive to customers and operations no matter when it actually occurs. From this perspective it is preferable to have a brief—but significant—disruption and transition completely to IPv6. This plan is akin to the idea that it is better to just rip the bandage off and have a moment of pain than removing it slowly in an attempt to reduce the pain.

The counterpoint to this argument is that we must look at the situation pragmatically with the goal of maintaining business continuity, growth, and customer satisfaction.

IPv4 Exhaustion

The impending IPv4 address exhaustion^[4] and the problems it will create has been the topic of much discussion in many different areas of the Internet community. The need to deploy IPv6 has figured prominently in the discussion, because it is the proper long-term solution. However, the unfortunate reality is that deploying IPv6 is a parallel activity to any work that provides continuity to the existing IPv4 network in order to keep it operational and able to grow to meet demand. As an Internet community, we are not where we need to be in terms of critical mass of our IPv6 deployments, in terms of either available, deployed equipment that supports IPv6 fully or applications that are able to use IPv6 when it is available.

IPv6 deployment is a requirement, but most ISPs do not have control over all variables affecting IPv6 deployment, and they have limited influence on progress outside of their network boundaries. This reality is especially true with residential services, where customers often purchase IP-enabled hardware directly from retailers to connect to their home networks. Consumers generally do not care about whether a device supports IPv4 or IPv6, so they do not make purchasing decisions based on such features. Customers should not be required to be technology experts in order to get their devices to work properly for their intended use. Customers generally are not interested in their ISP dictating the equipment that they may use in their home, and they do not like being told that they must replace “obsolete” gear, especially if they purchased it recently. The service provider sells “Internet” service, so customers expect their “Internet” devices to work—period. As a result, if an ISP wants to continue to grow, that ISP must continue to offer IPv4 services until the existing equipment without IPv6 support ages out of the network and is replaced.

The IETF recently released a *Best Current Practice* (BCP) document^[5] that provides some guidance for implementers that support for IPv6 on “IP-capable” devices is going to be a necessity, and the *Consumer Electronics Association* (CEA) now has a working group on IPv6 Transition^[6]. In conjunction with events like *World IPv6 Launch*^[7], there are near-constant improvements in the availability of IPv6-capable hardware, software, access, and services. The result of this situation should be that critical mass of IPv6 deployment will happen soon and reduce reliance on IPv4 and IPv4 life-extension technologies.

Because of the costs, operational complexities, performance concerns, and effects on customers that most IPv4 life-extension technologies create, service providers should focus on reaching IPv6 critical mass in essential areas.

When IPv6 has become sufficiently ubiquitous, the need for IPv4 life-extension technologies will be reduced along with the scale of deployments. Because a lot of the costs of deploying IPv4 life-extension technologies are initial costs, there is some truth to the argument that after they are deployed they are unlikely to disappear anytime soon. Why would a carrier invest significant time and money in deploying something only to pull it back out a short time later? Therefore the best method to reduce the cost of *Carrier-Grade NAT* (CGN) deployment is to work to deploy less of it.

ISPs are different when it comes to their expectations for growth, and their IPv4 addressing reserves or consumption rates differ accordingly. Some have areas of their internal network where they can make changes and reclaim globally unique IPv4 addresses for reuse to support customers, some have addresses that can be reclaimed via auditing and improved efficiency of allocation, and still others have already undertaken many of these projects and do not have much address space left to reclaim. Further, new IPv4 address availability as a combination of policies and demand may be different for each *Regional Internet Registry* (RIR). To summarize, the need for IPv4 address life-extension technologies is different on each network. The costs of deploying, the complexity of supporting, and the growth rate all figure into how widely service providers will have to deploy one or more technologies to extend their remaining IPv4 resources.

NAT444

Network Address Translation (NAT)^[28, 30] is already widely used for translating one IPv4 address to another, usually to provide separation or address sharing between a private network with multiple hosts and a public network or the Internet. In the context of IPv4 and IPv6 transition, these types of NAT are commonly referred to as *NAT44*, because they translate between IPv4 and IPv4 (vs. IPv4 to IPv6, IPv6 to IPv6, etc.). There is a proposed extension to NAT intended to preserve even more IPv4 resources. This proposal is called *Carrier-Grade NAT* (CGN)^[8]. The “Carrier Grade” in the name originates from the position of the NAT within the topology. Instead of NAT between a private and public network at the edge of a single network such as a home or business office, CGN is implemented inside of an ISP’s network and serves many customers simultaneously. These CGN implementations are typically scaled to handle thousands of simultaneous customer endpoints, often resulting in millions of simultaneous sessions. The RFC^[8] does not advocate the use of CGN; it describes how an ISP forced to deploy CGN can use it during IPv6 transition.

This sort of implementation addresses the need for an individual, globally unique IPv4 address for each of the ISP’s customers by allowing the ISPs to allocate each customer an IPv4 address that may not be globally unique and employ NAT to give them access to resources on the IPv4 Internet.

This sharing often allows ISPs to see oversubscription of public IPv4 addresses anywhere from 2:1 to more than 10,000:1 based on the type of applications behind the NAT and their simultaneous application layer port allocations and session counts. Most commonly, a CGN is used in conjunction with a local NAT on the customer's home network, creating two layers of NAT to traverse between the home network and the Internet. This model is commonly referred to as NAT444, because there is a translation layer between three sets of IPv4 addresses end to end.

A known problem with NAT is that it makes end-to-end communication and visibility between hosts more difficult, because it essentially hides hosts behind address translation. Because NAT is so common (nearly every home network and many commercial networks use NAT), networking applications have adapted so that they can discover the presence of a NAT and then change their behavior in order to maintain communications in the presence of NATs. However, the addition of this second layer of NAT often interferes with those workarounds, and undesirable or unpredictable results may occur^[9].

Over time it is likely that applications will again adapt to the impediments created by multiple layers of NAT, but it is not possible to anticipate and correct every potential problem that may be generated by adding this second layer of NAT. This reality should serve as a warning to those who provide services over an Internet connection: IPv6 support is extremely important. IPv6 is important because CGN means that ISP-controlled equipment will be actively involved in the path between content or application providers and their end users, making that relationship reliant on the service provider and the service provider's CGN vendor to an extent that was not necessary in the past. If the CGN implementation breaks something, it not only reflects on the CGN vendor and the service provider, it also reflects poorly on the relationship between the end customer and the service that that customer is using—and may cause that customer to form a negative opinion of the brand itself.

In other words, if a consumer uses an Internet-enabled application on a new Brand X smart TV and it does not work well, regardless of whether it is a problem with the CGN, the service provider, or something else entirely, the consumer may form the opinion and share via an online review that, "Brand X's TVs are ok, unless you try to use any of their fancy new features. I would not buy one if I were you, because Company X clearly does not know what it is doing." CGN represents a potentially significant increase in the amount of testing that must be done, especially in implementations that are uncommon, such as small, corner-case deployments, and closed architectures. Although using IPv6 is dependent on support at the client, the content or application provider, and the ISPs in between, if this support is present, it allows the content or application provider and client to bypass the service provider's CGN machinery—as well as any IPv4 NAT that may be present—and have a true end-to-end connection. This scenario restores control over the user experience back to the brand, and allows the ISP to resume supplying bit carriage.

IPv4 Addressing Requirements

Independent of the potential connectivity problems that NAT444 may create, it generates additional problems for the implementing ISP because of its need for IPv4 addresses. Because the CGN requires two sets of addresses—one for the inside (private) network and one for the outside (public) network—the ISP must identify address ranges to use for both. In order for its customers to be able to reach the Internet, the external pool must use globally unique IPv4 addresses. The number of addresses required will depend on the implementation of CGN, its scale profile, the topology of the network (how many hosts are behind each CGN instance), and the usage profile of the customer traffic. If the service provider has few or no available globally unique IPv4 addresses, it will have to either make changes in its network in order to reclaim addresses from elsewhere or make a request for a new allocation from its RIR^[29].

However, depending on the number of addresses that the RIR has available and its policies for justification, it may not be possible to obtain sufficient address space with this method. For example, in the Asia-Pacific region, the austerity policies in place mean that no matter how many IPv4 addresses they might have been able to justify using previous rules, most requesters are eligible for only a few hundred IPv4 addresses as their final allocation ever^[10]. This situation then requires the ISP to source IPv4 addresses via the IPv4 address transfer market^[11], adding additional cost to an already expensive deployment. In fact, if the service provider must source addresses via the transfer market, it may be more cost-effective to simply obtain more addresses and continue with business as usual without deploying CGN at all.

Internal Pool: Private Addressing Alternatives

When addresses are sourced for the public address pool, the service provider must also identify a pool of private addresses that is large enough for the provider to allocate one to each customer behind the CGN. Depending on the size and scale of the CGN, and how much the service provider is willing to segment and separate different sections of its network, this number could be a large block of addresses, perhaps even a /8 or more.

The most obvious choice might be to simply use address ranges reserved for private network use^[12], because there is a /8, a /12, and a /16 available for this purpose. However, this address space has some drawbacks. First, because of the prevalence of RFC 1918 addressing within most enterprise networks, there is a significant chance that the chosen address blocks may conflict with existing use of RFC 1918 space for management systems and other internal resources. Depending on the size of the CGN implementation, it may be necessary to instantiate multiple segments of the network where the entirety of RFC 1918 space is used, and in order for those segments to talk to one another or to talk to devices with conflicting numbering, significant additional complexity is required.

On the customer side, remote workers could experience problems where the address that they have been assigned is in a block that is already in use on their company's enterprise network, meaning that it may cause problems connecting to those hosts via a *Virtual Private Network* (VPN), or problems accessing some of the resources from the remote network. It may be possible to change the address assigned to the end user in an attempt to eliminate this conflict, but this approach is not necessarily scalable because it likely requires manual intervention in an automated address-assignment system, and there are limits to the number of times that a change of address can "fix" this problem without creating a problem for another user.

The other problem with the use of RFC 1918 space in the CGN is that it may conflict with the address space used by the customer's local network and NAT. For example, if a customer has a local network numbered out of **192.168.1.0/24** and the customer's router is allocated the external address of **192.168.1.85**, the router may fail to function properly because it has the same address range on both the internal and external interface. It may be possible through analysis to identify and carefully allocate addresses so that the portions of RFC 1918 commonly used by default in home gateway devices are not allocated. However, anecdotal evidence^[13] suggests that because of the wide variety of devices and implementations available—plus the fact that many users reconfigure their networks to use a different IP address range than the default configuration of the device—there simply may not be enough RFC 1918 addresses not in use to make this option viable.

"Squat" Space

Another alternative is to unofficially reuse one or more portions of the existing range of allocated globally unique IPv4 addresses as private addresses. In a network that does not talk directly to the Internet, such as a private network or VPN, the existing allocations of IPv4 space do not have any meaning, and so it is not strictly necessary to stick to RFC 1918 address space for numbering resources that are only internally accessible. Reuse of allocated IPv4 addresses has the benefit of not conflicting with in-use RFC 1918 addresses, but comes with its own set of problems. If the provider's own space is reused, the provider must carefully separate the private use from the public use to avoid conflicts, and managing this overlap may require additional complexity such as the use of VPNs as a method to separate the networks. The more common method is to reuse a block of addresses that is not currently allocated to the network using them; in other words, squatting on "someone else's" address space. Usually providers select space to use in this manner based on a low likelihood that either the owner will begin announcing the space on the global Internet or the users behind that network will need to connect to the users behind the ISP's NAT.

This method requires extreme care. The service provider must ensure that the routes for those prefixes are not inadvertently leaked to the global Internet, because such a leak could potentially cause a route-hijack denial-of-service attack, albeit an unintentional one. This method is even more risky if the ISP has one or more partners who have connections into the private portion of its network, because it may not have complete control of the announcement boundaries. Certainly there are safeguards such as tagging the announcements with *Border Gateway Protocol* (BGP) communities such as no-advertise or no-export^[14], but these solutions are not always practical, and they are not completely fail-safe. Depending on the chosen address space, the effects could be significant based on the true owner of that space—no service provider really wants to risk a public relations nightmare because it inadvertently caused an outage affecting the critical infrastructure of a large government agency or multinational corporation whose space it “borrowed” and then leaked to the Internet.

As a result of the IPv4 transfer market, it is quite likely that some of the address blocks that are not visible on the global Internet today and that some consider “safer” to squat on may end up being transferred to another party who plans to begin using them on the public Internet, and potentially requiring those squatting on the space to renumber to a different address block. ISPs can mitigate this risk somewhat by selecting multiple candidate blocks that are all preconfigured in the network such that it is relatively straightforward to make a rapid change from one block to another if the current block in use suddenly becomes unacceptable. Many ISPs use this method today, but because of the risks, it cannot be considered a real solution to the problem. Further, because it essentially encourages large service providers to violate the spirit—if not the letter—of the very policies that govern IP address allocation and use, standards bodies such as the IETF or policy organizations like RIRs cannot officially recommend such a solution.

Class E Addresses

A final alternative is to repurpose the reserved space in **240/4**^[2] and make it available for this use. There have been several failed attempts to repurpose this reserved space within the IETF in the past few years^[15, 16]. The primary challenge with this alternative is that because the Class E space has been reserved for many years, many networking implementations are explicitly configured to reject this address space as invalid. Getting this problem fixed in software, and more importantly, getting those software upgrades deployed widely, may require a similar level of effort to that which is required to deploy IPv6, and deploying IPv6 would be a more effective use of the resources required to implement software and hardware changes.

Even in situations like a CGN where more of the implementation is under central control, this solution would be attractive only to a service provider that owns and operates the *Customer Premises Equipment* (CPE) routers for all of its customers such that it could work with a small number of vendors to get software patches to enable use of this space. Therefore this solution is also too limited in applicability to be seen as a general solution that a body like the IETF could recommend.

Shared Addresses

Although the solutions previously discussed may be acceptable in some applications, the risks and deficiencies make it necessary for other applications to find another source for the IP address blocks to be used on the private side of a CGN. It is possible to use "public" (globally unique) IPv4 addresses on the private side as well, but the challenges to obtaining additional public IPv4 addresses that were discussed previously are exacerbated by the even larger number of addresses required, so this solution is far from practical. Additionally, expecting each service provider that implements CGN to obtain its own address space for its inside pools would end up using a significant amount of the remaining IPv4 resources in a way that does not necessarily require globally unique addresses. However, because each service provider has different needs, growth rates, and applications, it is unclear that simply expecting each service provider to request space from the RIRs for its internal CGN pools would create a doomsday scenario where a few networks would use up all of the remaining available IPv4 space in a short time. Because CGN creates additional costs and complexity to implement and support, and could be viewed as "second-class" IPv4 service, most service providers are not likely to implement it across the entire network and all tiers of customers, instead preferring to implement it only as widely as absolutely necessary.

Service providers could choose to implement it only for net new customers (that is, growth above turnover); they could choose to implement it only in certain markets or for certain types of service where it is less likely to cause support problems and adversely affect the service. All of these things reduce the number of addresses that may be needed for the interior CGN address pool. Nevertheless, using globally unique addresses in an application that does not require unique addresses is not a good use of a very limited resource. That is why the idea of having a shared and reserved block of addresses specifically for use as an interior (private) pool on a CGN keeps resurfacing.

One alternative to formally reserving a shared transition space was to have a third party request a block of sufficient size from one or more of the RIRs and then make it available for use as a shared block by anyone who wishes to do so.

Given the “last /8” policies in effect at each of the RIRs, it would likely be quite difficult to justify sufficient space to be useful, and the cost involved in receiving and maintaining such a delegation would likely be prohibitive. There would also be challenges addressing potential abuse concerns.

Reserving a block via the standard IETF/IANA process meant that IETF would have a chance to document the problems and recommend best practices that must be considered when implementing something that uses this shared space. This policy would help to ensure that service providers and implementers are aware of these guidelines and recommendations. For example, many implementations make certain assumptions about address scope based on the address itself, such as assuming that RFC 1918 addresses are locally scoped, and then adapt their behavior accordingly. With things like squat space or an unofficially shared CGN space, implementers would not know that this space should be treated in a specific way, and the result may be more network breakage. The officially declared shared space must still wait for implementers to make changes to their products, and that may not always happen, but the chances are still better than if it had been done in an unofficial manner.

As you can probably see, this problem does not have a clear-cut and straightforward solution, and this situation has led to vigorous discussion within the standards and policy bodies that have discussed it. The next section gives a brief history of the activity in those bodies that ultimately led to the space being allocated.

Some History

Shared transition space proposals have been controversial each time a variant of the idea has come up for discussion. As IPv4 exhaustion became a reality and IPv6 deployment continued to lag, more people realized that IPv4 life-extension technologies such as CGN may be a necessary evil. When people saw CGN as a likely response to the gap between IPv4 exhaustion and wide IPv6 support, they began to understand the need for the shared transition space, and thus support for allocating that space has gradually grown.

Although variants of this discussion may be much older than the items discussed in the following paragraphs, this article focuses specifically on the history of the idea to allocate shared address space specifically for CGN. There was an unsuccessful proposal in 2005^[17] to update RFC 1918 with an additional three /8s, but this proposal was not specifically focused on CGNs, unlike some of the other proposals. The most recent set of proposals regarding shared CGN space first came up in the APNIC Policy *Special Interest Group* (SIG) in early 2008, where Policy Proposal 058 was discussed. APNIC members abandoned the proposal and recommended that the authors take the idea to the IETF, because that is the body that typically directs IANA to reserve IP address blocks for special uses such as this one^[18].

This recommendation resulted in a pair of Internet drafts^[19, 20], hereafter referred to as **shirasaki** in late 2008. The draft originally requested four /8s, with a minimum size of a /12, but subsequent revisions of the draft revised the request to only one /10. The draft never gained much traction within the IETF, but the authors continued to update it to keep the discussion going. In mid-2010, a second IETF draft^[21] was published, requesting that a full /8 be reserved for this purpose. It contained references to the **shirasaki** drafts, but provided additional justification and noted that a /10 may not be enough addresses for many of the large service providers.

The draft went through several revisions in the following months, eventually being replaced by a different draft^[22], hereafter referred to as **draft-weil**, which reduced the /8 requested down to a /10. Attendees of the IETF 79 meeting in Beijing, China, discussed the draft across two different working groups. People expressed strong opinions both in support of and in opposition to the idea, but the draft did not achieve clear consensus. With the future of the draft unclear, one of its authors submitted policy proposal 127 to the *American Registry for Internet Numbers* (ARIN)^[23]. The ARIN *Advisory Council* (AC) accepted this policy proposal as draft policy 2011-5^[24] in early 2011, and vigorously discussed it with participants at the ARIN XXVII public policy meeting and with members of the mailing list. At the conclusion of the discussion, the ARIN AC recommended the policy to the ARIN board for adoption.

This discussion took on additional urgency because during this time the IANA officially announced that it had exhausted the free pool of IPv4 addresses and delegated the last of the /8s to the RIRs in accordance with policy^[4]. The side effect of this exhaustion meant that it was no longer possible for IETF to direct IANA to reserve space unless IANA was directed to repurpose an existing reservation, because it had no unreserved address blocks of sufficient size to meet the request. Therefore, the IETF and one or more of the RIRs would have to work in concert to make a suitable IPv4 address block available, instead of it being solely under IETF's purview. ARIN staff reached out to the IETF's *Internet Architecture Board* (IAB) for guidance, because by strict interpretation^[25], ARIN was not authorized to make this allocation by itself. IAB reaffirmed this interpretation, and recommended that the matter be brought back to the IETF for (re)consideration^[26]. With this guidance, the authors revised **draft-weil-shared-transition-space-request** and reintroduced it for discussion. For a period of time, the document was split into two, with most of the long-form discussion of pros and cons being moved to a second draft^[27].

As of the publication date of this article, the secondary draft has expired without progressing, but most of the important information contained there was incorporated back into **draft-weil**. The document was not adopted by any IETF Working Group. Instead, an IETF Area Director sponsored it as an individual submission.

It went through its first IETF “Last Call” to gauge consensus and receive comments in August 2011. The subsequent discussion, revisions, and secondary last calls (October 2011 and January 2012) generated hundreds of messages on the IETF discussion list and a total of 12 versions of the document before it was approved for publication in February 2012.

The reason why the debate on this shared transition space was so spirited can be traced to a few critical concerns. First, although consensus-based RFCs documenting CGN^[8] were already approved, this draft allocating space specifically to facilitate its deployment became a referendum within the IETF on whether NAT444/CGN should even be used. If you believed that NAT444 and CGN were bad ideas, it was likely that you would also be against a shared transition space. From that perspective, shared transition address space provided a more complete solution to a problem that had been created by a “Bad Idea” that should not have been allowed to proceed in the first place. There was also resistance to what was deemed “waste” of the limited remaining blocks of IPv4 addresses to solve a problem that not everyone agreed was a real or important problem. Also, although IETF participants do not speak for their companies per se, this proposal had consistent support from numerous individuals employed by large residential broadband providers. As a result, some saw it as those service providers looking for a way to bail themselves out of a problem that they created by not deploying IPv6 rapidly enough to avoid having to use CGN. On the converse side of the argument, those in favor saw CGN as a largely foregone conclusion, and saw this proposal as simply a practical solution to a real problem.

The *Internet Engineering Steering Group* (IESG) ultimately sent a note to the IETF discussion list acknowledging the difficulty of coming to a decision on this matter and noting that some explanatory text would be added to RFC 6598:

“Colleagues,

The IESG has observed very rough consensus in favor of the allocation proposed in **draft-weil-shared-transition-space-request**. Therefore, the IESG will approve the draft. In order to acknowledge dissenting opinions and clarify the IETF position regarding IPv6, the IESG will attach the following note:

“A number of operators have expressed a need for the special purpose IPv4 address allocation described by this document. During deliberations, the IETF community demonstrated very rough consensus in favor of the allocation.

While operational expedients, including the special purpose address allocation described in this document, may help solve a short-term operational problem, the IESG and the IETF remain committed to the deployment of IPv6.”

In many ways, the final decision came down to the difference between theory and practice in the IETF's desire to make the Internet work better. Theoretically, making a CGN easier to implement has the potential to make the Internet work much more poorly, and could be seen as rewarding bad behavior (failing to deploy and support IPv6 in a timely fashion). However, in practice, making CGN harder to implement causes unnecessary pain and effort for operators and potentially for users, while having little or no effect on IPv6 deployment. Approving this shared transition space avoids the appearance that IETF is trying to punish operators or users for perceived past "sins" and helps to reinforce the idea that IETF is responsive to operational concerns and therefore still relevant to the operator community. It is unlikely that the result of this decision will have much bearing on an operator's plan for how widely, when, where, or even if it will deploy CGNs, and this article makes no such recommendations. However, I will reiterate that IPv6 is the long-term solution, and that the smallest CGN deployment possible will make for a less complex and less expensive network for the continued support of traditional IPv4 devices.

Acknowledgements

Special thanks to Ole Jacobsen for suggesting that I write this article, to Kirk Erichsen and Jason Weil for their review and comments, and to all involved in the discussion of the referenced IETF drafts and ARIN policy for giving me plenty to write about!

References

- [1] Victor Kuarsingh, Chris Donley, Jason Weil, Marla Azinger, and Christopher Liljenstolpe, "IANA-Reserved IPv4 Prefix for Shared Address Space," RFC 6598, April 2012.
- [2] IANA Address Assignments:
<http://www.iana.org/assignments/ipv4-address-space/ipv4-address-space.txt>
- [3] RIR Policies triggered by IPv4 Depletion:
ARIN:
https://www.arin.net/resources/request/ipv4_countdown.html
RIPE:
<http://www.ripe.net/internet-coordination/ipv4-exhaustion/reaching-the-last-8>
APNIC:
<http://www.apnic.net/community/ipv4-exhaustion/ipv4-exhaustion-details>
LACNIC:
<http://www.lacnic.net/en/politicas/manual11.html>
AFRINIC:
<http://www.afrinic.net/docs/policies/AFPUB-2010-v4-005-draft-05.htm>

- [4] Number Resource Organization (NRO), "Free Pool of IPv4 Address Space Depleted," February 2011,
<http://www.nro.net/news/ipv4-free-pool-depleted>
- [5] Chris Donley, Christopher Liljenstolpe, Wesley George, and Lee Howard, "IPv6 Support Required for All IP-Capable Nodes," RFC 6540, April 2012.
- [6] Consumer Electronics Association IPv6 Working Group,
http://www.ce.org/Press/CurrentNews/press_release_detail.asp?id=12139
- [7] World IPv6 Launch, <http://www.worldipv6launch.org/>
- [8] Sheng Jiang, Brian Carpenter, and Dayong Guo, "An Incremental Carrier-Grade NAT (CGN) for IPv6 Transition," RFC 6264, June 2011.
- [9] Chris Donley, Lee Howard, and Victor Kuarsingh, "Assessing the Impact of Carrier-Grade NAT on Network Applications," Internet Draft, work in progress, November 2011,
[draft-donley-nat444-impacts-03](#)
- [10] APNIC, "Policies for IPv4 address space management in the Asia Pacific region,"
<http://www.apnic.net/policy/add-manage-policy#9.10>
- [11] ARIN, "Understanding the IPv4 Transfer Market,"
https://www.arin.net/resources/transfers/transfer_market.html
- [12] Daniel Karrenberg, Yakov Rekhter, Eliot Lear, and Geert Jan de Groot, "Address Allocation for Private Internets," RFC 1918, February 1996.
- [13] Akiro Kato, A sampling of RFC 1918 IP address usage in Japan,
<http://www.ietf.org/mail-archive/web/v6ops/current/msg06187.html>
- [14] Paul Traina, "BGP Communities Attribute," RFC 1997, August 1996.
- [15] Vince Fuller, "Reclassifying 240/4 as usable unicast address space," Internet Draft, work in progress, March 2008,
[draft-fuller-240space-02](#)
- [16] Paul Wilson, George Michaelson, and Geoff Huston, "Redesignation of 240/4 from 'Future Use' to 'Private Use,'" Internet Draft, work in progress, September 2008,
[draft-wilson-class-e-02](#)

- [17] Tony Hain, "Expanded Address Allocation for Private Internets," Internet Draft, work in progress, February 2005,
draft-hain-1918bis-01
- [18] Shirou Niinobe, Takeshi Tomochika, Jiro Yamaguchi, Dai Nishino, Hiroyuki Ashida, Akira Nakagawa, and Toshiyuki Hosaka, "Proposal to create IPv4 shared use address space among LIRs," prop-058, January 2008,
<http://www.apnic.net/policy/proposals/prop-058>
- [19] Jiro Yamaguchi, Yasuhiro Shirasaki, Shin Miyakawa, Akira Nakagawa, and Hiroyuki Ashida, "NAT444 addressing models," Internet Draft, work in progress, January 2012,
draft-shirasaki-nat444-isp-shared-addr-07
- [20] Ikuhei Yamagata, Shin Miyakawa, Akira Nakagawa, Jiro Yamaguchi, and Hiroyuki Ashida, "ISP Shared Address," Internet Draft, work in progress, January 2012,
draft-shirasaki-isp-shared-addr-07
- [21] Jason Weil, Victor Kuarsingh, and Chris Donley, "IANA Reserved IPv4 Prefix for IPv6 Transition," Internet Draft, work in progress, September 2010,
draft-weil-opsawg-provider-address-space-02
- [22] Victor Kuarsingh, Chris Donley, Jason Weil, Marla Azinger, and Christopher Liljenstolpe, "IANA-Reserved IPv4 Prefix for Shared Address Space," Internet Draft, work in progress, February 2012. (Became RFC 6598^[1]),
draft-weil-shared-transition-space-request-15
- [23] ARIN Public Policy Mailing List (PPML), "Shared Transition Space for IPv4 Address Extension," ARIN-prop-127,
<http://lists.arin.net/pipermail/arin-ppml/2011-January/019278.html>
- [24] "Shared Transition Space for IPv4 Address Extension," ARIN policy 2011-5,
https://www.arin.net/policy/proposals/2011_5.html
- [25] Brian Carpenter, Fred Baker, and Michael Roberts, "Memorandum of Understanding Concerning the Technical Work of the Internet Assigned Numbers Authority," RFC 2860, June 2000.
- [26] Internet Architecture Board, "Response to ARIN's request for guidance regarding Draft Policy ARIN-2011-5,"
<http://www.iab.org/documents/correspondence-reports-documents/2011-2/response-to-arins-request-for-guidance-regarding-draft-policy-arin-2011-5/>

- [27] Stan Barber, Owen Delong, Chris Grundemann, Victor Kuarsingh, and Benson Schliesser, "ARIN Draft Policy 2011-5: Shared Transition Space," Internet Draft, work in progress, September 2011,
draft-bdgks-arin-shared-transition-space-03
- [28] Geoff Huston, "Anatomy: Inside Network Address Translators," *The Internet Protocol Journal*, Volume 7, No. 3, September 2004.
- [29] Daniel Karrenberg, Gerard Ross, Paul Wilson, and Leslie Nobile, "Development of the Regional Internet Registry System," *The Internet Protocol Journal*, Volume 4, No. 4, December 2001.
- [30] Geoff Huston, "NAT++: Address Sharing in IPv4," *The Internet Protocol Journal*, Volume 13, No. 2, June 2010.
- [31] *The Internet Protocol Journal*, Volume 14, No. 1, March 2011. This issue of IPJ is entirely devoted to the topic of IPv4 address depletion and IPv6 transition.
- [32] Michelle Cotton and Leo Vegoda, "Special Use IPv4 Addresses," May 2012, **draft-vegoda-cotton-rfc5735bis-02**

WESLEY GEORGE has been working in IP networking for approximately 13 years, across operations, engineering and capacity planning, architecture, and design in large wired and wireless networks. He has been heavily involved in IPv6 evangelism and deployment for a surprisingly long time. He has been an active participant in IETF for 5 years, including serving as former co-chair of the IPv6 Renumbering (6renum) working group and current co-chair of the sunset4 working group. He was active in ARIN's policy development process during the time that the policy discussed in this article was being addressed. He currently works for Time Warner Cable, but this article represents his views alone, and should not be mistaken for his current employer's official stance on anything. Wes can be reached via twitter (@wesgeorge) or via wesley.george@twcable.com

December in Dubai: Number Misuse, WCIT, and ITRs

by Geoff Huston, APNIC

In November 1988, telephone companies from 178 nations sent their respective government representatives to the *World Administrative Telegraph and Telephone Conference* (WATTC) in Melbourne, Australia. At the time the generally cosy relationships between governments and their monopoly telephone companies often made it extremely difficult to see the difference between the government's representatives and those of the telephone company. The group resolved to agree to the rather grandly titled *International Telecommunication Regulations* (ITRs)^[1].

At this meeting the companies' national representatives agreed to a set of additional regulations that supplemented the binding regulations of the *International Telecommunication Convention*. The goals of these regulations were rather grand; they aspired to promote the "harmonious development and efficient operation of technical facilities, as well as the efficiency, usefulness and availability to the public of international telecommunication services." More practically, these ITRs defined the general principles for the provision and operation of international telephony services among signatories to the ITRs.

At that time the Internet was little more than a somewhat obscure experiment in advanced data communication protocols undertaken by a small number of researchers in North America and to a far smaller extent in Europe. However, since 1988 the Internet—and the world in which the Internet has flourished—has changed dramatically. If we view the rise of the Internet over the past 25 years as a product of an appropriately liberalized international regulatory regime as much as it was a product of the titanic shifts in computing and communications technologies that also occurred over this period, then we can make the case that the Internet of today is a product of these ITRs. And what a prodigious product it has been!

In Dubai, between the 3rd and 14th of December 2012, the nations of the world will convene at the 2012 *World Conference on International Telecommunications* (WCIT)^[2], and they intend to use this conference to review these 25-year-old ITRs and consider some proposed changes to this regulatory framework that underlie international telecommunications.

At the moment the international meeting cycle is ramping up to consider what aspects of the ITRs should be altered, what should stay the same, and what should be dropped. After all, much has happened in the past 25 years, and an argument could be made that the ITRs should be amended to better reflect today's world.

But the world is not exactly aligned at the moment about what should and what should not be folded into a new set of international regulatory obligations.

Some countries appear to be advocating for some quite specific measures to be added to the ITR to address what for them are characterized as otherwise unresolvable operational problems. Others are advocating a more general approach to have the ITRs explicitly embrace the Internet and fold references to the Internet in every place where specific carriage and service delivery technologies are referenced in the ITRs. It is when these two approaches intersect that the situation gets interesting.

In order to illustrate some of the underlying tensions that exist in this activity, I would like to take a specific example of a proposed amendment to the ITRs and consider in in terms of the broader context of telephony and the Internet.

The proposal I want to examine here concerns the topic that has been called "number misuse." In telephony this term referred to an operating practice where a call to a dialled number is not routed to the destination subscriber who is located at that called number, but instead the call is re-routed to a different destination.

What we see in the "Number Misuse" proposal for a revision of the ITRs is an attempt to fold the concepts of "number misuse" and the Internet together, with a result that some countries want the ITRs to explicitly take on the concept of "IP Address and Routing Misuse" within the framework of national obligations through common regulatory action within the same scope as the telephony called number misuse. If successful, this effort would result in a regulatory obligation for governments to take necessary actions to investigate and prosecute such instances of so-called "number misuse." The intended scope of such enforcement of such obligations would encompass not only the telephone network but also the Internet. Surely we all desire a global public communications network that operates with integrity, and surely we would want to see countries take the necessary actions to ensure that it happens. So why is this idea not exactly the best idea to appear in the ITR negotiation process so far?

Let's look at the motivations behind number misuse in the world of telephone carriers and telephone services, and then look at how it could conceivably map in to the world of the Internet.

To understand the telephone world and where this problem of number misuse is coming from, it may be useful to understand a little of how money circulates in the phone world.

Telephony: Sender Pays

In many ways the telephone leaned heavily on the telegraph service for its service model, which, in turn, leaned on the postal service, establishing a provenance for the telephone service model that stretched back over some centuries to at least the 1680s and London's Penny Post, if not earlier.

The postal service model that gained ascendancy over the preceding centuries was one in which the original sender of the letter paid for the entire service of letter delivery. If the postal service that received the letter in the first place needed to use the services of a different postal service to complete the delivery, neither the sender nor the intended recipient were aware of it. The postal services were meant to divide the money received from the sender to deliver the letter, and apportion it between themselves to compensate each service provider for undertaking its part in the delivery of the letter.

The telephone service, for the most part, operates in a very similar fashion. The caller pays for the entire cost of the call, and the called party pays nothing.

When both the caller and the called party are connected to the same carrier, the process is straightforward. The carrier charges the caller for the cost of the call and, presumably, some small (often not so small) margin for profit.

However, when we apply the same model to, say, international phone calls, the model is not so simple. The common desire on the part of the telephone operators was to preserve the same simple model: the caller pays. Now in this case the caller pays the presumably higher price of establishing a voice circuit from a carrier in one country in one part of the world to another carrier in another country in another part of the world. But now the caller's carrier should not keep all the revenue associated with the call. The other end, the *terminating carrier*, has also incurred costs in servicing this call. The arrangement that the telephone industry developed was the concept of "intercarrier call accounting financial settlements."

To explain this concept it may be useful to introduce the unit of a *call minute*, which is commonly used as a means of measuring a telephone call. What carriers establish between themselves on a bilateral basis is the intercarrier settlement cost per call minute of a telephone call that originates in one carrier and is terminated by the other carrier.

Now if both carriers can establish a value of a call-minute settlement rate where in both directions the call-minute termination costs roughly equate to the call-minute settlement rate, then in theory, at any rate, neither party is relatively advantaged over the other, irrespective of whether the callers are predominately located in one carrier or in the other carrier. In theory, such an arrangement should be financially neutral to both carriers.

However, although in theory practice and theory should align, in practice it rarely happens. What happened in the telephone case was that we saw some carriers set a call-minute call-termination settlement rate that was well above cost, while at the same time set its international call tariffs such that outbound calls were prohibitively expensive for local subscribers.

The result was that the local customers of these carriers found it cheaper to request that the other party call them—the desired outcome. The local carrier then generated income not by charging local subscribers but by revenue generated as an outcome of the call accounting settlement payments that were generated by the net imbalance of called versus calling call minutes.

Carriers all over the world played this game. For example, in France in the early 1990s it was some 5–10 times more expensive to call a U.S. number from France than it was to make a call between the same two numbers in the other direction. If you add in a further consideration, namely that in the 1980s many carriers were part of the public administration and were in effect government-operated national monopolies whose profits contributed to national revenue, then you get an outcome that is described in *Opinion No. 1* of the 1989 ITRs, under the heading “Special Telecommunication Arrangements,” namely: “...considering further that, for many Members, revenues from international telecommunications are vital for their administrations.”

Telephony Special Services and Number Misuse

It is often said that the only really major innovation in more than a century of the telephone service was the fax. Perhaps that is a little too unkind, but innovations in the delivered services industry were few and far between. However, there were many innovations that are important to this story, and the ones that are relevant here are *number redirect* and the so-called *premium* services.

The premium services attracted a higher call cost, and the carrier conventionally split the revenue from the service with the called service. These services traditionally included weather forecasts, sports results, new headlines (until the Internet became all but completely ubiquitous and decimated these services!), and so on. They also attracted the sex industry. However, in many countries such services were not permitted, so a conventional premium service was not an option for this industry.

As ever, we are naturally inventive, and some folks came up with a clever solution to use number redirect to redirect the call to this otherwise not-permitted premium service to another country. As part of this redirection, the premium service provider needed to reach an agreement with the new home carrier of the call-termination point to divide the international call accounting revenue provided by callers to this service between the carrier and the service provider. Not only did this arrangement effectively circumvent local regulations relating to locally provided premium services, it also leveraged off the international call accounting arrangements to the benefit of the premium service provider as well as the terminating carrier.

We may be inventive, but all too often we are greedy as well. The next step was to circumvent any arrangement with the destination carrier and redirect the call to an entirely different carrier.

One of the side effects of deregulation of the telephone industry in many countries was that in place of a single carrier that would receive all incoming international calls for a given country code there were numerous carriers that were ostensibly competing for these incoming calls. Instead of routing calls based solely on the dialled country code, carriers now could route calls based on number blocks within the country code, and use different transit routes based on number-block rules. What if a premium service provider took a number block from a country code and specified that all incoming calls were to be routed by a third-party carrier? That all sounds innocent enough, but what if this third party did not actually route the calls through to the country in question, but instead terminated the calls and still charged the calling carrier the international call accounting settlement rate? No doubt the service provider has gotten a better deal, so the service provider is happy, and the carrier that terminates the call is receiving a portion of the call settlement rate, so the terminating carrier is happy. But happiness is not universal here. The carrier in the called country code is getting nothing from this arrangement, even though its country call code is being used for these premium service calls. From the carrier's perspective it is being defrauded of what it might claim is legitimate international call accounting revenue through the "misuse" of the number block drawn from its country code.

If the country-code carrier could discover this unauthorized number-block diversion, then presumably it could withdraw the number block and stop the international call diversion. Unfortunately it does not always work. The carrier can withdraw the number block, but at times—and under perhaps somewhat shady circumstances—the premium service provider, and potentially the transit carriers, might still be able to convince local carriers that the number-block diversion is still legitimate. Although the country-code carrier might see the problem, the carrier's ability to enforce carriers in other countries to respect its authority regarding the use of number blocks drawn from its country code is not always clear. At times the carrier is effectively powerless to enforce a remedy.

And the scheme can be further refined. Why even enter into any form of discussion with the international carrier for a number block? Why not pick one or more of the more obscure national country codes, generate some number blocks from these codes, and then get a cooperative transit carrier to enter a number-block diversion request into the local carrier? The number block is perhaps drawn from a country code that already makes extensive use of third-party transit arrangements, the local carrier may not question the request, and the carriers in the countries from which the number blocks have been drawn may not have the resources to even detect that this event has occurred.

At this point we have arrived at the situation that is motivating some of the proposals to augment the ITRs in this round of negotiation. The position of the nations that have been highlighting this problem as being an important problem in the world of international telephony is that the unauthorized use of phone numbers drawn from their E.164^[3] telephone number block is, in their eyes, a case of "number misuse."

The reason why they want to identify this situation and write it into the ITRs at this time is that they would like to involve governments in the role of enforcers of conformance with the conventions of management of telephone country codes. It appears that they would like to obligate governments to adopt a policy, as a common convention, that calls made to a country's country code be directed such that the call request is sent to an authorized carrier located in the country, and to ensure that all authorized carriers essentially honor the integrity of the country codes of all other countries that use the E.164 country-code number plan.

It is also reasonable to ascribe the motivation for this measure as one that is intended to ameliorate the inexorable revenue leakage of the former rich money tap of international call accounting settlement payments. I am not sure that the various antics of the international premium service market are the true intended target of this measure. I suspect that the intended targets of this proposed regulatory measure are those carriers that have devised other methods to honor the intentions of their callers when they make an international phone call, and make the phone of the dialled number ring, yet at the same time bypass the traditional call accounting arrangements. Already *Voice over IP* (VoIP) trunking is commonplace, where the call is mapped into a VoIP call, and one way to bypass the conventional call accounting measures is to use a VoIP trunk to enter the dialled country, and then pass the call back into the *Public Switched Telephone Network* (PSTN) as a locally originated call, terminating it on the originally dialled number. The call is then subject to domestic intercarrier call-termination tariffs, which are generally far lower than their international counterparts.

The Internet and services such as Skype are exerting massive downward pressure on what carriers can charge for conventional phone services without encouraging all remaining customers to use Internet-based services. In an effort to retain some level of market share, it is now evidently more commonplace for carriers themselves to embrace IP-based approaches and bypass these imposed intercarrier international settlement charges. For many countries in the developing world, however, this shift represents a twofold financial blow. Not only are they seeing their foreign-sourced revenue stream disappear at the same rate as the call-termination minutes of conventional telephony vaporise, but they are also seeing this revenue stream being replaced by growing IP traffic volumes that represent a net cost to the national economy.

It should come as no surprise to see some countries attempt to advocate an international regulatory response that is intended to reverse this development, and restore the role of the international telephone network as a means of structural flow of monies from the business sector from the richer economies to the consolidated revenue stream of those poorer economies.

Internet Number Misuse

In and of itself, the previous discussion is by no means a novel discussion for the telephone world, and the tensions exposed by the continual erosion of the traditional telephone business through the onslaught of new technology is not at all surprising.

What is perhaps a bit surprising are the recent moves within the ITR preparatory activities that see numerous national delegations advocating pulling Internet addressing and routing into the same category of telephone-number regulation and also fold these factors into this matter of number misuse in a manner that would apply to both E.164 numbers and IP addresses.

Now some things do not readily translate from telephony to the Internet: there is no “National IP Address Plan” as a counterpart to the E.164 number plan, because the IP address plan is aligned to networks, as distinct from countries. However, you could take a broad view and find some form of mapping from the proposed recommendations regarding the use of E.164 networks to IP addresses. It would appear that the application of the proposals regarding number misuse would see a regulation to the effect that IP packets should be routed to the destination address specified in the packet, and not rerouted and terminated elsewhere. Surely this scenario describes part of the way the Internet works in any case. For the network to actually function, packets need to be passed to their addressed destination. Or so you would think.

And that is indeed what happens much of the time within the Internet. But by no means all of the time. As part of the normal course of operation of IP networks, many operators deploy equipment that intercepts packets and forms a synthetic response using the address of the intended destination. And many national administrations either operate—or mandate the operation of—equipment that inspects packets in transit and discards packets addressed to certain number blocks.

What is going on? Why do network operators regularly “misuse” IP addresses by deliberately intercepting packets and generating a synthetic response?

Packet Diversion

The most prevalent reason is the use of proxies, and, in particular, web proxies. These devices sit “on the wire” and intercept web fetches and cache the downloaded data.

When another user requests the same URL, the proxy uses the cached version of the content rather than forwarding the request on to the original site. This caching is by no means unusual: it is typical for web browsers to cache the most recently visited webpages and when the user returns to the page, the local cached copy is used rather than re-performing the download. For the browser and the network operator the rationale for this form of "address misuse" is the same: it is both a desire to improve performance for the end user and a desire to increase the efficiency of the network by reducing the data volumes being shifted across the transit links. So the outcomes are, on the whole, positive outcomes; users see improved performance and potentially lower costs for the service, using an interception technique that is generally transparent.

Is the deployment of a web proxy an instance of fraud?

Here is where another critical difference between the Internet and the telephone world comes into play. In the Internet the sender does not "pay all the way" to get a packet from its source to its intended destination. In general, every IP packet could be thought of as being partially funded by both the sender and the receiver.

The user who generated the packet pays for an *Internet Service Provider* (ISP) service, and the ISP may, in turn, purchase transit services from another ISP, and so on for sequenced transit services. However, at a peering exchange point, or within a provider network, the sender's money runs out. The packet is not unfunded, however, for at this point the receiver's services take over, and the packet transits a path that is funded by the receiver's ISP's transit services, and there to the receiver's ISP and there to the receiver.

If a packet is diverted to a proxy, then who wins and who loses? Can we make the case that a party in this situation is being cheated?

As long as the proxy is a faithful proxy, then the user wins, insofar as the user experiences improved performance and the benefits of a more efficient network while still seeing precisely the same content. And the content provider wins, insofar as the content is delivered to the user without the incremental cost of packet handling at the content site. And the network service providers win, in so far as the amount of network traffic is reduced while the revenue levels remain constant. In this case there is no end-to-end service payment on the part of the user that would trigger an intercarrier settlement payment, so it is difficult to make the case that this action necessarily damages any party involved in the network transaction.

Given the widespread deployment of these proxy caching devices across the entire Internet, the beneficial outcomes of improved performance and network efficiency, and the option for content providers to use techniques that in effect mark content as not cacheable, it is extremely challenging to sustain a case that the use of proxies is a case of address misuse.

So the use of traffic diversion and intercepting proxies in the Internet is not generally regarded as an example of intentional fraud or even an accepted case of address misuse. It is just what we do today in the Internet.

Packet Interception

What about the deliberate interception and discarding of packets in flight? Surely this case is one of "misuse" of IP addresses?

That is a very hard case to make when you consider that such actions are exactly how firewalls work, and almost every network uses firewalls in some manner or other. The action of a firewall is to intercept all packets, and discard those that match some predetermined set of rules relating to acceptable and unacceptable packets.

Many users run firewalls that deliberately block all incoming connection requests unless they match quite specific rules.

Many ISPs run firewalls that deliberately block access to ISPs' services from users who are not direct customers of the ISP.

Many countries have content regulations that block access to certain content, enforced either through government-operated facilities or through obligations imposed through the conditions associated with the carrier license within that country. The country I live in, Australia, imposes such constraints on its carriers for certain types of content, as does China through its much-reported national firewall facilities.

Users, service providers and carriers, and governments all use various forms of packet interception. Are we all guilty of number misuse? Should we support changes to the ITRs to obligate governments to stop this practice completely?

Aside from many other motivations for firewalls, security is a continuing concern in the Internet, and there is little doubt that although firewalls have not eradicated all forms of toxic traffic and associated abuse and attack, they are an important part of a larger story about securing the Internet. Irrespective of the various views that are expressed at a national level about censorship, intellectual property rights, and the position of common carriers and users, it seems counterintuitive to me that we would want to obligate governments to pull down our firewalls and filters as a necessary consequence of a revised set of ITRs.

Number "Misuse"

What this example illustrates is that the two networks—the traditional telephone network and the Internet—operate in very distinct and different ways. It not only encompasses differences between circuit and packet switching, but also reaches into the differences in the concepts of a network transaction, differences in the tariff structures, and, critically, differences in the way in which financial settlements are undertaken between service providers on the Internet.

Consider what could readily be acknowledged as an operating practice that defrauds operators in the world of telephony and negatively affects the services provided to telephone subscriber—that same practice in the Internet can result in positive outcomes used to enhance performance, reduce costs, and improve the operational efficiency of the service delivered to end users.

This case of attempting to regulate “number misuse” illustrates the fact that to take a stance of “one size fits all” when considering the topic of international regulation of telecommunications is a stance that has considerable risks of generating outcomes that are entirely inappropriate when translating a particular situation from telephony to the Internet.

WCIT and the ITRs—Where to Go from Here?

The international call accounting arrangements used by the telephone world, and the use of structurally embedded imbalances in call accounting settlement rates, are still major factors in the ITR discussions. This accounting imbalance is sanctioned in the resolutions of the 1988 World Administrative Telegraph and Telephone Conference, where *Resolution 3*, concerning the apportionment of revenue, provided for structural cross-subsidization of the developing world through asymmetric fixing of call accounting rates between the so-called developed and developing economies.

But in an increasing commercial world of telecommunications, where it is no longer a relatively exclusive collection of publicly funded monopolies that were an integral part of public utility service providers that in effect were an instrument of national governments, pushing the onus of an international developmental agenda onto an increasingly privatized commercial activity has been a less-than-comfortable fit. Private operators see this situation in a more dispassionate light as a business cost input, and seek to find ways to minimize this cost in order to improve the competitive positions of their businesses.

However, the changes in this industry over the past 25 years are so much larger than even this significant broad-scale shift in the onus of capital injection and operation from the public to the private sector. At the same time, we are seeing an even more fundamental shift in technology foundations, from circuits to packets with the introduction of the Internet into the picture. This shift has brought about profound shifts in the engineering of communications infrastructure and, as we have seen, it also has triggered profound shifts in the pricing of the consumer service, shifting from transactional pricing to a “connection rental” model where packet transit costs are bundled into the service. This bundling, in turn, has led to profound shifts in the manner in which money moves between the network operators themselves.

And perhaps of even greater and more lasting significance in this industry is the decoupling of carriage and content. We have now seen the rise of highly valuable content-centric enterprises that have business models that rely on a ubiquitous and abundant underlying communications infrastructure but are not financially beholden to the infrastructure operators. They have been able to forge direct relationships with consumers without having to deal with any form of mediation or brokerage imposed by carriage providers. The current values of these content enterprises dwarf the residual value of the carriage service sector, and the outlook for this sector is one of continuing shift in value away from carriage service providers and into the areas of content-based services.

Given the sheer scale of these changes in this industry over the past quarter century, it seems to me that the view that you can simply fold the Internet transparently into the current framework of the ITRs by the prolific insertion of "and the Internet" into the text of the regulations is simply not viable.

Packets are not circuits, and the mechanisms used to engineer packet networks are entirely different from those used with the circuit switches that supported traditional telephony services. This difference encompasses far more than engineering. The ways in which users pay for services differ, and this shift in the retail tariff structure of the Internet service implies a forced change in the way in which carriers interact to support a cohesive framework of network interconnection. The concept of a "call" really has no direct counterpart in the Internet. To extend this thought further into the area of "call accounting" and "caller pays" is again an extension that does not clearly map into the Internet. So when the existing ITRs refer to intercarrier call accounting financial settlements, there is no clear translation of such a concept into the Internet. When we extend this intercarrier interconnection framework into structural imbalances in call accounting settlement rates, and extend this framework further into the concepts of number misuse, all forms of connection between traditional telephony and the Internet are completely lost.

However, this conclusion should not imply that the ITRs are now an historic relic, completely overtaken by comprehensive shifts in both the technology and service models of today's global communications network. Irrespective of the fine level of detail in these 25-year-old documents, the ideals behind the ITRs are indeed worthy ideals, and they should not be discarded lightly.

Ultimately, what we are dealing with here is the role of individual nation states with respect to a public communications service for the entire world. In setting forth a framework for supporting an efficient, effective, and capable global communications system, the obligations stated in the current ITRs relating to the promotion of international telecommunications services, and the endeavours to make such services generally available to the public, all remain thoroughly worthwhile objectives.

The concept that widely respected technology standards are critical to worldwide technical interoperability of any telecommunications service is also an important aspect, and again the recognition of this factor in the ITRs is a worthwhile consideration.

But, as we both review the changes of the past quarter century and try to peer into what may emerge over the next quarter century, perhaps less is best in this area of regulatory measures.

Rather than seeking to explicitly add various regulations that attempt to address specific incidents of number misuse, and instead of making rather clumsy efforts to include the Internet into the already detailed provisions relating to intercarrier settlement models of the increasingly historic traditional telephone network, perhaps the best set of ITRs we could have for tomorrow's world are national obligations that support a lightweight common regulatory framework.

This framework should be both more minimal with respect to describing or relying on particular technologies and service frameworks and more encompassing in scope in stating the overall objectives and common aspirations all nations share in supporting this unique, incredibly valuable common resource of a common communications service that truly embraces the entire world.

Postscript: "It's All Just Telecoms"

I received a comment soon after I wrote an early draft article that I thought would provide some further insight to the WCIT process, so here is the comment and some further thoughts on the topic:

The comment was in the form of a report from a preparatory meeting for WCIT earlier in 2012. Evidently there is a mood within certain parts of the ITR drafting process to simply say: "The ITRs should apply to the Internet in full, because the Internet is nothing more than a telecom service and should be treated that way."

In one sense it is true that the Internet is nothing more than a telecommunications service, but in the same way that the post, radio, television, and of course the telephone are also all just telecommunications services. But the nature of the particular service has many consequences, and the attempt to lump telephony and the Internet into the same form of regulatory handling is at best a somewhat misguided effort.

I truly wonder if, more than a century ago, the counterparts of today's government delegates, in a meeting of that august body, the *Universal Postal Union* (UPU), would have argued that a telephone conversation was just an exchange of letters without the artifice of paper, and that the telephone was indeed just a part of the postal service, because it is just "a communications service."

Indeed I am pretty sure their counterparts did precisely that, and for the next 80 years or more in many countries the Postmaster General operated the telephone service, and operated the wireless spectrum administration and regulated radio and television broadcasts, as well as operating the national postal service, the telegraph service, and telex services, all because "it's all just communications."

But, ultimately we changed this paradigm. We created distinct entities to administer different communications media and services because it is actually not "all just communications"—nor is it "all just telecoms." Effective regulatory handling of these different communications mechanisms, using distinct forms of investment and finances, and at times entirely distinct regulatory frameworks and often distinct organizations and associated participatory arrangements, allows us to realize the true potential of these various services and do so efficiently and effectively. This recognition of a need for distinction in the regulatory frameworks for various services avoids the unfortunate situation of the stultifying dead hand of history misapplying one form of regulation on an entirely distinct and very different medium.

I suspect the best thing the postal folks, in the form of the UPU, ever did was to tell the telephone folks "hail and farewell" and let them get on with their role using an organization specifically designed to meet their collective needs in supporting telephony.

It may be well and truly time for the telephone folks, in the form of the *International Telecommunications Union* (ITU), to come to a similar arrangement in its dealings with the Internet!

Disclaimer

These views do not necessarily represent the views or positions of the Asia Pacific Network Information Centre.

Further Reading:

- [1] The current International Telecommunication Regulations (1988):
http://www.itu.int/dms_pub/itu-t/oth/3F/01/T3F010000010001PDFF.pdf
- [2] World Conference on International Telecommunications (WCIT-12),
<http://www.itu.int/en/wcit-12/Pages/default.aspx>
- [3] Geoff Huston, "ENUM—Mapping the E.164 Number Space into the DNS," *The Internet Protocol Journal*, Volume 5, No. 2, June 2002.

GEOFF HUSTON, B.Sc., M.Sc., is the Chief Scientist at APNIC, the Regional Internet Registry serving the Asia Pacific region. He has been closely involved with the development of the Internet for many years, particularly within Australia, where he was responsible for the initial build of the Internet within the Australian academic and research sector. He is author of numerous Internet-related books, and was a member of the Internet Architecture Board from 1999 until 2005. He served on the Board of Trustees of the Internet Society from 1992 until 2001.
E-mail: gih@apnic.net

Behind the Curtain: IP Fast Reroute

by Russ White, Verisign

The field of network and protocol engineering has three watchwords: *faster*, *bigger*, and *cheaper*. Although we all know the joke about choosing two out of the three, the reality of networking is that we have been doing all three for years—and it doesn't look like there is any time on the horizon when we will not be doing all three.

In that spirit, *IP Fast Reroute* addresses all three of these watchwords. Fast—you are probably thinking—is obvious, but what about bigger and cheaper? Fast Reroute provides the network designer with some trade-offs in the space of redundancy through additional backup links against deploying protocol changes, and network stretch against the size of a failure domain, so you can—in theory—build larger, less-redundant failure domains with Fast Reroute than without.

But to understand these effects, we need to go behind the curtain, understanding Fast Reroute as more than a few configuration options. This article first looks at the motivation behind IP Fast Reroute, and then discusses four different techniques, or stages, in the Fast Reroute story.

What Is Your Motivation?

To really discuss network speed, we need to be able to define how fast “fast” really is. In the 1980s, a network was fast if it could converge in 90 seconds or less (the longest time the *Routing Information Protocol* [RIP] could take to converge). As we moved into more advanced Distance-Vector and Link State protocols (*Enhanced Interior Gateway Routing Protocol* [EIGRP], *Open Shortest Path First* [OSPF], and *Intermediate System-to-Intermediate System* [IS-IS]), 5-second convergence became the norm. We learned to tweak timers to get to convergence times faster than 1 second.

But what if we need convergence that is faster than less than 1 second? What if we need to converge so fast that the only packets lost are either in flight or in a buffer waiting to be serialized onto the link? And what if we need to be able to handle a large number of prefixes with minimal network disruption due to link or device failures?

IP Fast Reroute techniques come into play in this situation.

Preinstalled Backup Paths

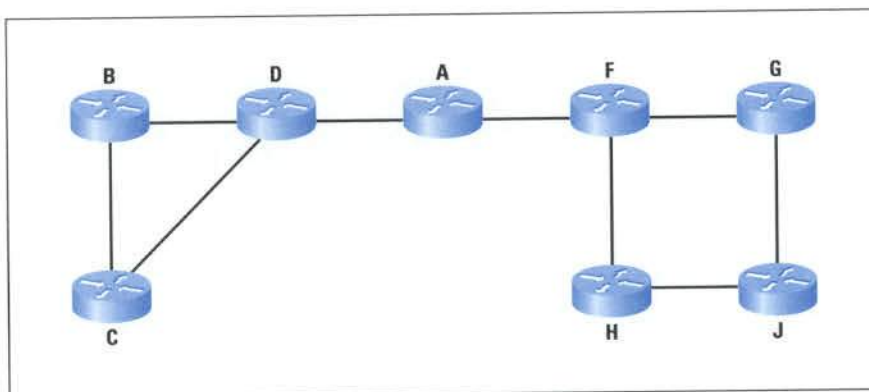
Although it is often sold as a Fast Reroute technique, *preinstalled backup paths* really are not; rather they support other Fast Reroute techniques at the protocol level. If the protocol has calculated a loop-free path that is an alternate to the current best path, this alternate path can be installed in the forwarding table so it is readily available for use in case the best path fails.

This solution does provide immediate failover at the hardware level, but the alternate path must be calculated to be installed. How is this alternate path computed?

Loop-Free Alternates

The first mechanism available for calculating an alternate path is with *Loop-Free Alternates*. To understand this mechanism, we must make a short detour into graph theory (or geometry, if you prefer). Use the following network as an example:

Figure 1: Network for Loop-Free Alternates



Assume:

- A is the destination.
- B's best path is through D to A.
- G's best path is through F to A.

What is the key to allowing B to forward traffic through C toward A if the $B \rightarrow D$ link fails? B must know the traffic it forwards to C (for A) will not be forwarded back to B itself. How can B know C will forward the traffic to D, rather than to B itself? By examining the metric at C toward A.

In EIGRP, B knows C's metric toward A because the routing protocol includes this information in the update. In a link state protocol (OSPF or IS-IS), B can calculate C's cost to A directly by running *Shortest Path First* from C's perspective (given B and C share the same link state database).

Loop-free alternates are simply calculating whether any given neighbor will forward traffic to any particular destination back to you, or on toward the destination. If a neighbor would forward the traffic on toward the destination, then it is a loop-free alternate.

Under what conditions would C forward traffic sent from B back to B? If C is using B as its best path (or one of its best paths) toward A.

What about G? If it forwards traffic to J toward A, will J return the traffic to G itself? In this four-hop ring, there are two possible configurations:

- J is using H as its best path. In this case, traffic forwarded by G to A through J will be correctly forwarded. Note, however, that in this case H cannot use J as an alternate path toward A, because any traffic H sends to A through A will loop back to H itself.
- J is using G as its best path. In this case, J can use G as a loop-free alternate, but G cannot use J as a loop-free alternate.

No matter how you work the metrics in the four-hop ring case, there will always be at least one device that does not have a loop-free alternate path to A.

Split Horizon and Loop-Free Alternates

If the concept of loop-free alternates is difficult to understand by considering the problem in this way, another useful way to look at the problem is through the distance-vector idea of *split horizon*. To review, the split horizon rule states:

Do not advertise a route to a destination toward a neighbor you are using to forward traffic to that same destination.

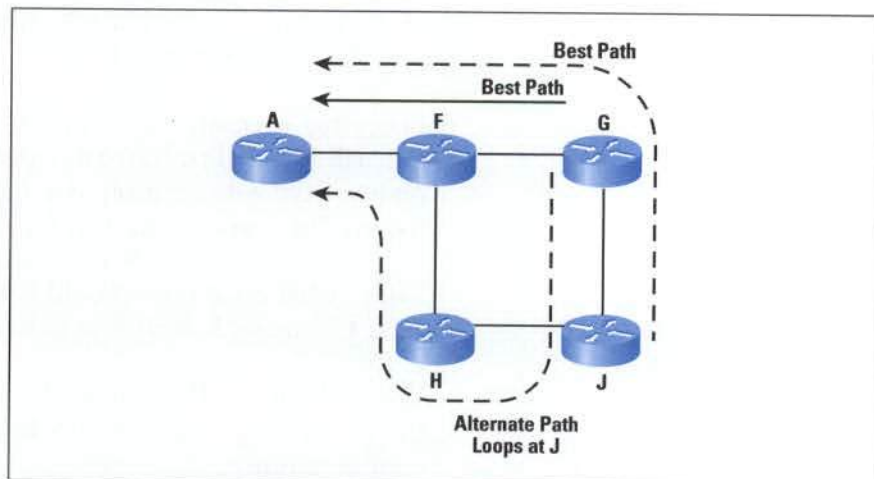
If C is forwarding traffic toward A to B, then C will not advertise A to B, meaning B will not even know about this alternate path, preventing a loop even if B's best path to A fails. If you always consider where a distance-vector protocol will split horizon, you will always be able to see where loop-free alternates will fail to provide an alternate path to any given destination.

Getting Around the Loops

If we want to design a system that will find every possible alternate path toward a given destination, rather than just finding those that are not normally taken out by split horizon anyway, what must we do? We need to find a way to route through a neighbor to some distant next hop without that neighbor actually forwarding the traffic back to the originating router.

To put this concept in more concrete terms, examine the following network as an example:

Figure 2: Alternate Path Loops



If G wants to use the path through J as an alternate path, then it must somehow figure out how to forward traffic to J without J returning the traffic to G itself. How can this process be done? G can tunnel the traffic through J to some device somewhere beyond J; therefore, every mechanism beyond loop-free alternates must use some form of tunneling to resolve the Fast Reroute problem. Calculating the point to which G needs to tunnel is the topic of the remaining mechanisms.

Not-Via

Even though we might be working with a link state protocol, it is easiest to understand *Not-via* in terms of a distance-vector protocol and split horizon. Not-via essentially begins with the observation that G does not have an alternate path to A through J in this case because J will not advertise such a route. *J is, in fact, using G as its best path toward A, so the path from G through J to A cannot be viable.*

The solution is not just simply having J advertise the route to A because traffic forwarded by G toward A through J will simply be looped back to G itself. So what is the solution?

In the case of Not-via, F advertises a route to itself through H only (not through G). This route will be advertised through H, then J, and finally to G. When G receives this route, it can determine that this path is an alternate path to A because its best path to A is normally through F. Any path that can reach F not through (or not via) its best path to F must, necessarily, be a loop-free alternate path to F. To reach A through F, however, G must tunnel to F directly, thereby avoiding the problem of J returning traffic destined to A back to G.

The address F advertises through H only is called “F Not-via G,” and that is why this system is called “Not-via.” This mechanism works in every topology (so long as an alternate path exists). The one downside to Not-via is that for each protected link or node, a new advertisement must be built and advertised through the network.

Disjoint Topologies

The problem of finding a next hop that passes over the split-horizon point can also be solved using the ability to form multiple disjoint topologies—multiple topologies that do not share the same links (or nodes, in some cases) to reach the same set of destinations. If this information sounds complex, that is because it is complex; a lot of hours and thought have gone into various systems to build and use multiple disjoint topologies within a single physical network. But there is a moderately simple way, referring back to Figure 2. In this network, G can take the following steps:

1. Remove the $G \rightarrow F$ link from its local database temporarily (just for this calculation).
2. Calculate the best path to F.
3. If an alternate path to F exists, mark this alternate path as a second topology.

4. If its path to F fails, place all traffic that would normally pass across $G \rightarrow F$ on this alternate topology.

It might not be obvious from this set of actions, but these actions will actually cause G to discover that it is, in fact, on a ring, and that it can place traffic on the opposite direction on this ring to get traffic to the same destination. Placing the traffic it would normally send to F via $G \rightarrow F$ on a separate topology overcomes the forwarding table at J, a process that would loop the traffic back to G itself. You could use a tunnel to F instead of a separate topology; tunnels are, in effect, a disjoint topology seen in a different way.

Conclusion

What advantage does IP Fast Reroute provide the network designer? The ability to reduce the amount of physical redundancy while maintaining the same actual level of redundancy in the network. Moving to Not-via or disjoint topology solutions removes the need to manually manage link costs as well, while adding only moderate complexity at the protocol level.

IP Fast Reroute is an interesting technology just on the edge of adoption that will be useful in campus, data center (through Layer 2 routing), and standard Layer 3 network designs.

For Further Reading

Work is currently active on the disjoint topology mechanism within the research community and the IETF; in particular, the following drafts will be of interest to anyone who wants to learn more:

- [1] Alia Atlas, Robert Kebler, Maciek Konstantynowicz, Andras Csaszar, Russ White, and Mike Shand, "An Architecture for IP/LDP Fast-Reroute Using Maximally Redundant Trees," Internet Draft, work in progress, October 2011,
draft-atlas-rtgwg-mrt-frr-architecture-01
- [2] Alia Atlas, Gabor Envedi, and Andras Csaszar, "Algorithms for Computing Maximally Redundant Trees for IP/LDP Fast-Reroute," Internet Draft, work in progress, March 2012,
draft-envedi-rtgwg-mrt-frr-algorithm-01
- [3] Stefano Previdi, Mike Shand, and Stewart Bryant, "IP Fast Reroute Using Not-via Addresses," Internet Draft, work in progress, December 2011,
draft-ietf-rtgwg-ipfrr-notvia-addresses-08
- [4] Clarence Filsfils and Pierre Francois, "LFA applicability in SP networks," Internet Draft, work in progress, January 2012,
draft-ietf-rtgwg-lfa-applicability-06

RUSS WHITE is a Principle Research Engineer at Verisign. He has co-authored numerous technical books, RFCs, and software patents. He focuses primarily on network complexity, network design, the space where routing and naming intersect, control-plane security, protocol design, protocol operation, and software-defined networks. E-mail: riwhite@verisign.com

Letters to the Editor

Ed.: We received several letters in response to the article "A Retrospective: Twenty-Five Years Ago," by Geoff Huston, published in the previous issue of this journal. Here is some of the feedback:

Hi Geoff,

Just wanted to show my appreciation for your nice article. As an ex-DEC who moved to WorldCom after my MSc in Computer Engineering & Telecoms with a Master's project on IP signaling over ATM, I can certainly relate to a large part (not all ;-)) of what you wrote.

I normally don't read such long articles, but had to make an exception as I kept interested until the end!

Thank you!

—Pedro Paiva, Etoy, Switzerland
`pedro.paiva@a3.epfl.ch`

Greetings Geoff,

I just wanted to let you know that I really enjoyed your recent article, "A Retrospective: Twenty-Five Years Ago," published in *The Internet Protocol Journal*. I lived through most of the history that you talked about as I came up through the telecom industry and then finished off my career at Cisco.

It certainly is interesting to reflect back on all the past controversy around network infrastructure design and how competing ideas and philosophies played out. (Talk about losers, remember *Switched Multi-megabit Data Service* (SMDS) driven by the *Regional Bell Operating Companies* (RBOCs)? While at Nortel, I remember once in a design review meeting that one of our BNR geeks put up a slide (overhead foil back then) that showed various network evolution scenarios. The last one was an "oh-by-the-way, there's this theory that the Internet could take over the world" (of network infrastructure). All the room snickered. Who's laughing now?

There was as much energy, maybe more, put into defending architectures based on market control as there was on technological elegance. Still, it is a fascinating and dynamic industry full of extremely smart people with clever ideas, and I enjoyed every minute of it.

I started at “the phone company” in the late 1960s and it has been quite a journey from relay-driven switches controlling tip and ring loops to the current *Multiprotocol Label Switching* (MPLS) backbone networks, terabit switching, and hitching rides on photons.

Thanks for your insight and for your well-written article.

Best regards,

—Marc Williams
willimarc@gmail.com

The author responds:

Hi Marc,

Thanks for your note and your recollections from some 25 years ago.

I recall SMDS as well. If I recall correctly, this was an invention coming out of a university in Western Australia. Elsewhere in the world it was marketed as a 34-Mbps product. In Australia it was marketed in 2-Mbps and 10-Mbps forms (evidently the telco thought that we primitive Aussies were not “ready” for any higher speed!). I was a customer of their 10-Mbps product, and experienced some disappointment when it became evident that 10 Mbps was a theoretical peak that was simply unachievable because the inline PCs that were used for packet accounting slowed the throughput of any SMDS link down to just 3 Mbps! So in Australia SMDS was largely killed by the telco and it was never really used for high-speed digital trunk services.

I experienced a similar reaction to the Internet in the late 1980s as you have observed, when, in response to suggesting that the universities were about to build a national IP network, many of the telco managers did the polite snicker performance and then suggested that we should “get with the times,” sign up as customers of their national ATM network, and leave the engineering to them. I’m glad the universities saw through it and supported me in persisting along the path to a national IP network. It was a strange moment some 6 years later when the same telco came knocking on our door to make an offer to buy the network from the universities because their own efforts to construct an IP product were simply getting nowhere at the time.

It has indeed been quite a journey, and I too have enjoyed every bit of it!

Kind regards,

—Geoff, Chief Scientist, APNIC
gih@apnic.net

Hello Geoff!

I haven't chuckled that much in years; what great memories. A few of my strong memories:

- Lack of documentation for new functions in software required an off-net test network and a Sniffer. The amount of hours spent figuring exactly what the function was doing or wasn't doing could fill an ocean. Absolutely my favorite activity and still is.
- I inherited a stat-mux system that was transporting ASCII terminals back to a centralized DEC terminal server arrangement. Hated it with a passion. One day, after a couple of beers, a light bulb came on that Ethernet is a stat-mux, so I bought a couple of Cisco AGS units, remotely installed a terminal server and an AGS, hauled it back to the other AGS in the central location and danced a jig, and then I started ripping out the old WAN stat-mux the following week.
- Anything relying on a token for timing is pure evil. You never know when you've engineered a TTL exhaust until it happens, and that can be based on Distance + Nodes or pure application coincidence. Ring resets are the devil's work. Token-based systems are not stat-muxs, but Ethernets are; that's why Ethernet survived and is the "last man standing."
- I totally agree with your comments surrounding the "cloud." I can remember that the distributed-versus-centralized fad has occurred at least four times over the past 25 years ...
- Z80: I built my first PC with a Z80; thank goodness for the peek-and-poke function!
- OEM would claim anything was portable as long as it had a carrying handle attached, even if it took two people to carry it.
- I fell in love with TCP/IP very early for the simple reason that it has the best of both worlds: a tightly coupled connection and connectionless protocol. It is much faster to troubleshoot or modify because IP requires a different expertise than TCP, and when you run across individuals who can work across the layers, hire them!

So, a lot of fond memories. I started out as a telemetry engineer on the Apollo project and I thought that was challenging and fulfilling. But, it doesn't hold a candle to the 1984–1995 period.

Oh, one other thing; I take umbrage to "...the annoying persistence of FORTRAN." That's the first language I learned back in the late '60s and I still have an active compiler on an old laptop that I still program on ... LOL!!

Keep attacking the certificate situation! The current situation is a disgrace, and I fully support the concept presented by Barnes: let's hurry it up!

Regards,

—Paul Dover
pdover@centeriem.com

The author responds:

Hi Paul,

Thanks for those recollections. I too spent a massive amount of time starting as a protocol analyzer, trying to make an IBM PC look enough like a Uniscope to allow file transfer between the PC and the Univac mainframe—no doubt it was a character-forming experience, but all I can say now is thank goodness for *tcpdump* and *wireshark*!

Thanks for your note—I truly appreciate the feedback!

Warm regards,

—Geoff, Chief Scientist, APNIC
gih@apnic.net

Dear Ole,

Congratulations on your 25-year anniversary!

You can tell how well people enjoy their professions by how great their products are, and yours is in the “excellent” category.

Regards,

—Paul Dover
pdover@centeriem.com

Ole,

Congratulations on your reaching a major milestone: 25 years of technology publishing! We are glad that you are continuing this service through *The Internet Protocol Journal* and look forward to many more years in this field.

Best,

—T. Sridhar
tsridhar@ieee.org

Call for Papers

The Internet Protocol Journal (IPJ) is published quarterly by Cisco Systems. The journal is not intended to promote any specific products or services, but rather is intended to serve as an informational and educational resource for engineering professionals involved in the design, development, and operation of public and private internets and intranets. The journal carries tutorial articles ("What is...?"), as well as implementation/operation articles ("How to..."). It provides readers with technology and standardization updates for all levels of the protocol stack and serves as a forum for discussion of all aspects of internetworking.

Topics include, but are not limited to:

- Access and infrastructure technologies such as: ISDN, Gigabit Ethernet, SONET, ATM, xDSL, cable, fiber optics, satellite, wireless, and dial systems
- Transport and interconnection functions such as: switching, routing, tunneling, protocol transition, multicast, and performance
- Network management, administration, and security issues, including: authentication, privacy, encryption, monitoring, firewalls, troubleshooting, and mapping
- Value-added systems and services such as: Virtual Private Networks, resource location, caching, client/server systems, distributed systems, network computing, and Quality of Service
- Application and end-user issues such as: e-mail, Web authoring, server technologies and systems, electronic commerce, and application management
- Legal, policy, and regulatory topics such as: copyright, content control, content liability, settlement charges, "modem tax," and trademark disputes in the context of internetworking

In addition to feature-length articles, IPJ contains standardization updates, overviews of leading and bleeding-edge technologies, book reviews, announcements, opinion columns, and letters to the Editor.

Cisco will pay a stipend of US\$1000 for published, feature-length articles. Author guidelines are available from Ole Jacobsen, the Editor and Publisher of IPJ, reachable via e-mail at ole@cisco.com

This publication is distributed on an "as-is" basis, without warranty of any kind either express or implied, including but not limited to the implied warranties of merchantability, fitness for a particular purpose, or non-infringement. This publication could contain technical inaccuracies or typographical errors. Later issues may modify or update information provided in this issue. Neither the publisher nor any contributor shall have any liability to any person for any loss or damage caused directly or indirectly by the information contained herein.



The Internet Protocol Journal, Cisco Systems
170 West Tasman Drive
San Jose, CA 95134-1706
USA

ADDRESS SERVICE REQUESTED

PRSRT STD
U.S. Postage
PAID
PERMIT No. 5187
SAN JOSE, CA

The Internet Protocol Journal

Ole J. Jacobsen, Editor and Publisher

Editorial Advisory Board

Dr. Vint Cerf, VP and Chief Internet Evangelist
Google Inc, USA

Dr. Jon Crowcroft, Marconi Professor of Communications Systems
University of Cambridge, England

David Farber
Distinguished Career Professor of Computer Science and Public Policy
Carnegie Mellon University, USA

Peter Löthberg, Network Architect
Stupi AB, Sweden

Dr. Jun Murai, General Chair Person, WIDE Project
Vice-President, Keio University
Professor, Faculty of Environmental Information
Keio University, Japan

Dr. Deepinder Sidhu, Professor, Computer Science &
Electrical Engineering, University of Maryland, Baltimore County
Director, Maryland Center for Telecommunications Research, USA

Pindar Wong, Chairman and President
Verifi Limited, Hong Kong

*The Internet Protocol Journal is
published quarterly by the
Chief Technology Office,
Cisco Systems, Inc.
www.cisco.com
Tel: +1 408 526-4000
E-mail: ipj@cisco.com*

*Copyright © 2012 Cisco Systems, Inc.
All rights reserved. Cisco, the Cisco
logo, and Cisco Systems are
trademarks or registered trademarks
of Cisco Systems, Inc. and/or its
affiliates in the United States and
certain other countries. All other
trademarks mentioned in this document
or Website are the property of their
respective owners.*

Printed in the USA on recycled paper.

