

The Internet Protocol Journal

December 2001

Volume 4, Number 4

*A Quarterly Technical Publication for
Internet and Intranet Professionals*

In This Issue

From the Editor	1
Scaling Inter-Domain Routing.....	2
Regional Internet Registries ...	17
Book Reviews	30
Letters to the Editor.....	34
Fragments	38
Call for Papers	39

FROM THE EDITOR

In a previous article entitled “Analyzing the Internet BGP Routing Table,” Geoff Huston examined many issues relating to the operation of today’s Internet. In this issue he goes a step further and suggests ways in which the fundamental routing architecture could be changed to solve problems related to routing-table growth. The article is called “Scaling Inter-Domain Routing—A View Forward.”

The IP address space is administered by three entities, namely APNIC, ARIN and RIPE NCC. Collectively referred to as the *Regional Internet Registries* (RIRs), these organizations are responsible for address allocation to their member organizations (typically national registries or large Internet Service Providers). Since the IPv4 address space is a limited resource, this allocation has to be done with care, while accounting for the needs of the address space consumers. We asked the RIRs for an overview of the work they perform. What we received was a joint effort that not only describes the RIR structure, but also gives some historical background on the evolution of IP addressing and routing.

We were pleased to receive a couple of Letters to the Editor recently, both in response to articles in our previous issue. This kind of feedback is most welcome and we encourage you to send your comments and suggestions to ipj@cisco.com

We’d like to remind you that all back issues of *The Internet Protocol Journal* can be downloaded from www.cisco.com/ipj. Click on “IPJ Issues” and you will be taken to the appropriate section.

By the time you read this, our online subscription system should be operational. You will find it at our Web site: www.cisco.com/ipj. Please let us know if you encounter any difficulties by sending e-mail to ipj@cisco.com

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

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

Scaling Inter-Domain Routing—A View Forward

by Geoff Huston, Telstra

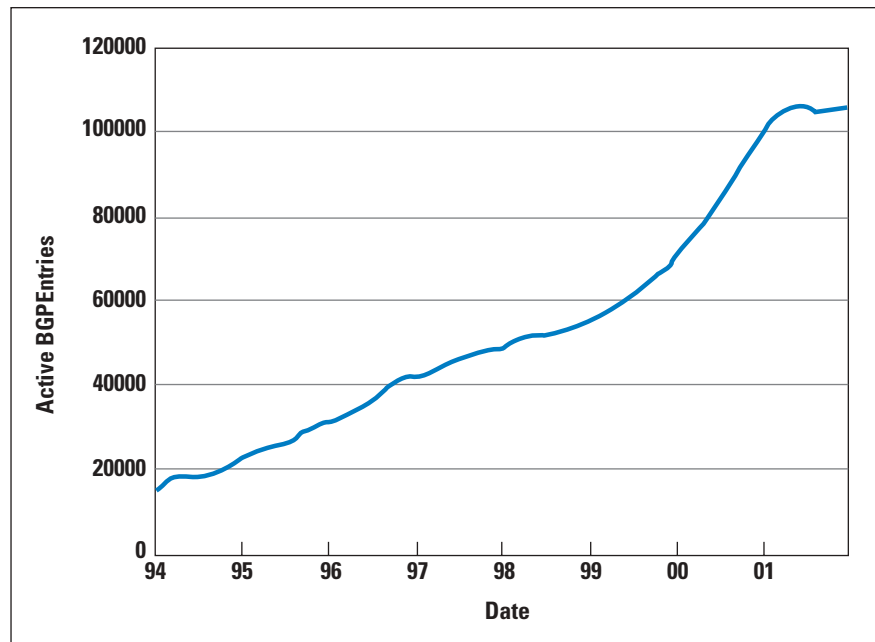
In the previous IPJ article, “Analyzing the Internet BGP Routing Table,” (Vol. 4, No. 1, March 2001) we looked at the characteristics of the growth of the routing table in recent years. The motivation for this work is to observe aspects of the Internet routing table in order to understand the evolving structure of the Internet and thereby attempt to predict some future requirements for routing technology for the Internet.

The conclusions drawn in the previous article included the observation that multihomed small networks appeared to be a major contributor to growth of the Internet routing system. It also observed that there was a trend toward a denser mesh of inter-Autonomous System connectivity within the Internet. At the same time there has been an increase of various forms of policy-based constraints imposed upon this connectivity mesh, probably associated with a desire to undertake various forms of inter-domain traffic engineering through manipulation of the flow of routing information.

Taken together, these observations indicate that numerous strong growth pressures are being exerted simultaneously on the inter-domain routing space. Not only is the network itself growing in size, but also the internal interconnectivity of the network is becoming more densely meshed. The routing systems that are used to maintain a description of the network connectivity are being confronted with having to manipulate smaller route objects that describe finer levels of network detail. This is coupled with lengthening lists of qualifying attributes that are associated with each route object. The question naturally arises as to whether the *Border Gateway Protocol* (BGP) and the platforms used to support BGP in the Internet today can continue to scale at a pace that matches the growth in demands that are being placed upon it.

The encouraging news is that there appears to be no immediate cause for concern regarding the capability of BGP to continue to support the load of routing the Internet. The processor and memory capacity in current router platforms is easily capable of supporting the load associated with various forms of operational deployment models, and the protocol itself is not in imminent danger of causing network failure through any internal limitation within the protocol itself. Also, numerous network operators have exercised a higher level of care as to how advertisements are passed into the Internet domain space and, as a result, the growth rates for the routing table over 2001 shows a significant slowdown over the rates of the previous two years (Figure 1).

Figure 1: BGP Table
Size 1994–2001



However, the observed trends in inter-domain routing of an increasingly detailed and highly qualified view of a more densely interconnected and still-growing network provide adequate grounds to examine the longer-term routing requirements. It is useful, therefore, to pose the question as to whether we can continue to make incremental changes to the BGP protocol and routing platforms, or whether the pace of growth will, at some point in time, mandate the adoption of a routing architecture that is better attuned to the evolving requirements of the Internet.

This article does not describe the operation of an existing protocol, nor does it describe any current operational practice. Instead it examines those aspects of inter-domain routing that are essential to today's Internet, and the approaches that may be of value when considering the evolution of the Internet inter-domain routing architecture. With this approach, the article illustrates one of the initial phases in any technology development effort—that of an examination of various requirements that could or should be addressed by the technology.

Attributes of an Inter-Domain Routing Architecture

Let's start by looking at those aspects of the inter-domain routing environment that could be considered a base set of attributes for any inter-domain routing protocol.

Accuracy

For a routing system to be of any value, it should accurately reflect the forwarding state of the network. Every routing point is required to have a consistent view of the routing system in order to avoid forwarding loops and black holes (points where there is no relevant forwarding information and the packet must be discarded). Local changes in underlying physical network, or changes in the policy configuration of the network at any point, should cause the routing system to compute a new distributed routing state that accurately reflects the changes.

This requirement for accuracy and consistency is not, strictly speaking, a requirement that every node in a routing system has global knowledge, nor a requirement that all nodes have precisely the same scope of information. In other words, a routing system that detects and avoids routing loops and inconsistent black holes does not necessarily need to use routing systems that rely on uniform distribution of global knowledge frameworks.

Scalability

Scalability can be expressed in many ways, including the number of routing entries, or prefixes, carried within the protocol, the number of discrete routing entities within the inter-domain routing space, the number of discrete connectivity policies associated with these routing entries, and the number of protocols supported by the protocol. Scalability also needs to encompass the dynamic nature of the network, including the number of routing updates per unit of time, time to converge to a coherent view of the connectivity of the network following changes, and the time taken for updates to routing information to be incorporated into the network forwarding state. In expressing this ongoing requirement for scalability in the routing architecture, there is an assumption that we will continue to see an Internet that is composed of a large number of providers, and that these providers will continue to increase the density of their interconnection.

The growth trends in the inter-domain routing space do not appear to have well-defined upper limits, so placing bounds on various aspects of the routing environment is impractical. The only practical way to describe this attribute is that it is essential to use a routing architecture that is scalable to a level well beyond the metrics of today's Internet.

In the absence of specific upper bounds to quantify this family of requirements, the best we conclude here is that at present we are working in an inter-domain environment that manipulates some 10^5 distinct routing entries, and at any single point of interconnection there may be of the order of 10^6 routing protocol elements being passed between routing domains. Experience in scaling transmission systems for the Internet indicates that an improvement of a single order of magnitude in the capacity of a technology has a relatively short useful lifetime. It would, therefore, be reasonable to consider that a useful attribute is to be able to operate in an environment that is between two to three orders of magnitude larger than today's system.

Policy Expressiveness

Routing protocols perform two basic tasks: first, determining if there is at least one viable path between one point in the network and another, and secondly, where there is more than one such path, determining the "best" such path to use. In the case of interior routing protocols, "best" is determined by the use of administratively assigned per-link metrics, and a "best" path is one that minimizes the sum of these link metrics.

In the case of the inter-domain routing protocols, no such uniformly interpreted metric exists, and “best” is expressed as a preference using network paths that yield an optimal price and performance outcome for each domain.

The underlying issue here is that the inter-domain routing system must straddle a collection of heterogeneous networks, and each network has a unique set of objectives and constraints that reflect the ingress, egress, and transit routing policies of a network. Ingress routing policies reflect how a network learns information, and which learned routes have precedence when selecting a routing entry from a set of equivalent routes. In a unicast environment, exercising control over how routes are learned by a domain has a direct influence over which paths are taken by traffic leaving the domain. Egress policies reflect how a domain announces routes to its adjacent neighbors. A domain may, for example, wish to announce a preferential route to a particular neighbor, or indicate a preference that the route not be forwarded beyond the adjacent neighbor. In a unicast environment, egress routing policies have a bearing on which paths are used for traffic to reach the domain. Transit routing policies control how the routes learned from an adjacent domain are advertised to other adjacent domains. If a domain is a transit provider for another domain, then a typical scenario for the transit provider would be to announce all learned routes to all other connected domains. For a multi-homed transit customer, routes learned from one transit provider would normally not be announced to any other transit provider.

This requirement for policy expressiveness implies that the inter-domain routing protocol should be able to attach various attributes to protocol objects, allowing a domain to communicate its preferences relating to handling of the route object to remote domains.

Robust Predictable Operational Characteristics

A routing system should operate in such a way that it achieves predictable outcomes. The inference here is that under identical initial conditions a routing system should always converge to the same routing state, and that with knowledge of the rules of operation of the protocol and the characteristics of the initial environment, an observer can predict what this state will be. Predictability also implies stability of the routing environment, such that a routing state should remain constant for as long as the environment itself remains constant.

The routing protocol should operate in a way that tends to damp propagation of dynamic changes to the routing system rather than amplify such changes. This implies that minor variations in the state of the network should not cause large-scale instability across the entire network while a new stable routing state is reached. Instead, routing changes should be propagated only as far as necessary to reach a new stable state, so that the global requirement for stability implies some degree of locality in the behavior of the system.

The routing system should have robust convergence properties. A change in the physical configuration or policy environment in any part of the network causes a distributed computation of the routing state. Convergence implies that this distributed computation reaches a conclusion at some point. The requirement for a robust convergence property implies that the distributed computation should always halt, that the halting point be reached quickly, and the system should avoid generating transitory incorrect intermediate routing states. The interpretation of “quickly” in this context is variable. Currently, this value for BGP convergence time is of the order of tens to hundreds of seconds. In order to support increasingly time-critical applications, there appears to be an emerging requirement to reduce the median convergence time for the inter-domain routing protocol to a small number of seconds.

Efficiency

The routing system should be efficient, in that the amount of network resources, in terms of bandwidth and processing capacity of the network switching elements, should not be disproportionately large. This is an area of trade-off in that the greater the amount of information passed within the routing system and the greater the frequency of such information exchanges, the greater the level of expectation that the routing system can continuously maintain an accurate view of the connectivity of the network, but at a cost of higher overhead. It is necessary to pass enough information across the system to allow each routing element to have a sufficiently accurate view of the network, yet ensure that the total routing overhead is low.

Evolving Requirements of Inter-Domain Routing

Layered on top of the base set of routing requirements listed above are a second set of requirements that can be seen as reflecting current directions in the deployed Internet, and are not necessarily well integrated into the existing routing architecture.

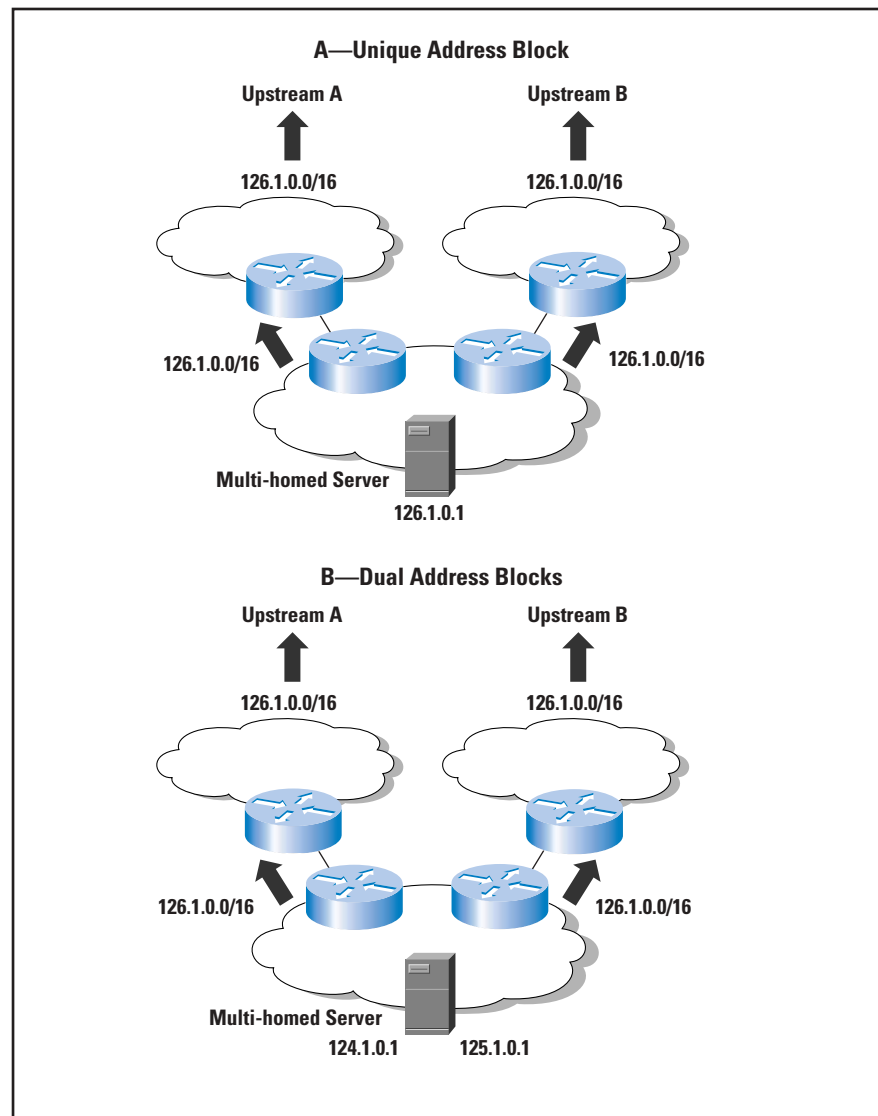
Multi-Homing of Edge Networks

Multi-homing refers to the practice of using more than one upstream transit provider. The common motivation for such a configuration is that if service from one transit provider fails, the customer can use the other provider as a means of service restoration. It may also allow some form of traffic balancing across multiple services. With careful use of route policies, the customer can direct traffic to each provider to minimize delay and loss, achieving some improved application performance.

The issue presented by multi-homing is that the multi-homed network is now not wholly contained within a service hierarchy of any particular provider. This implies that routing information describing reachability to the multi-homed customer cannot readily be aggregated into any single provider’s routing advertisements, and the usual outcome is that the multi-homed customer must independently announce its reachability to each transit provider, who in turn must propagate this information across the routing system.

The evolving requirement here is one that must be able to integrate the demands of an increasing use of multi-homing into the overall network design. Two basic forms of approach can be used here—one is to use a single address block across the customer network and announce this block to all transit providers as an unaggregatable routing advertisement into the inter-domain routing system, and the other is to use multiple address blocks drawn from each provider's address block, and use either host-based software or some form of dynamic address translation within the network in order to use a source address drawn from a particular provider's block for each network transaction (Figure 2). The second approach is not widely used, and for the immediate future the requirement for multi-homing is normally addressed by using unique address blocks for the multi-homed network that are not part of any provider's aggregated address blocks. The consequence of this is that widespread use of multi-homing as a means of service resiliency will continue to have an impact on the inter-domain routing system.

Figure 2: Routing Approaches to Multi-Homing



Inter-Domain Traffic Engineering

In an increasingly densely interconnected network, selecting and using just one path between two points is not an optimal outcome of a routing architecture. Of more importance is the ability to identify a larger set of viable paths between these points and distribute the associated traffic flows in such a way that each individual transaction uses a single path, but the total set of flows is distributed across the set of paths.

To achieve this outcome, more information must be placed into the routing system, allowing a route originator to describe the policy-based preferences of which sets of paths should be preferred for traffic destined to the route originator, allowing a transit service operator to add information regarding current preferences associated with using particular transit paths, and allowing the traffic originator the ability to use local traffic egress policies to reach the destination. These traffic engineering-related preferences are not necessarily represented by static values of routing attributes. One of the requirements of traffic engineering is to allow the network to dynamically respond to shifting traffic load patterns, and this implies that there is a component of dynamic information update that is associated with such traffic engineering-related aspects of the routing system.

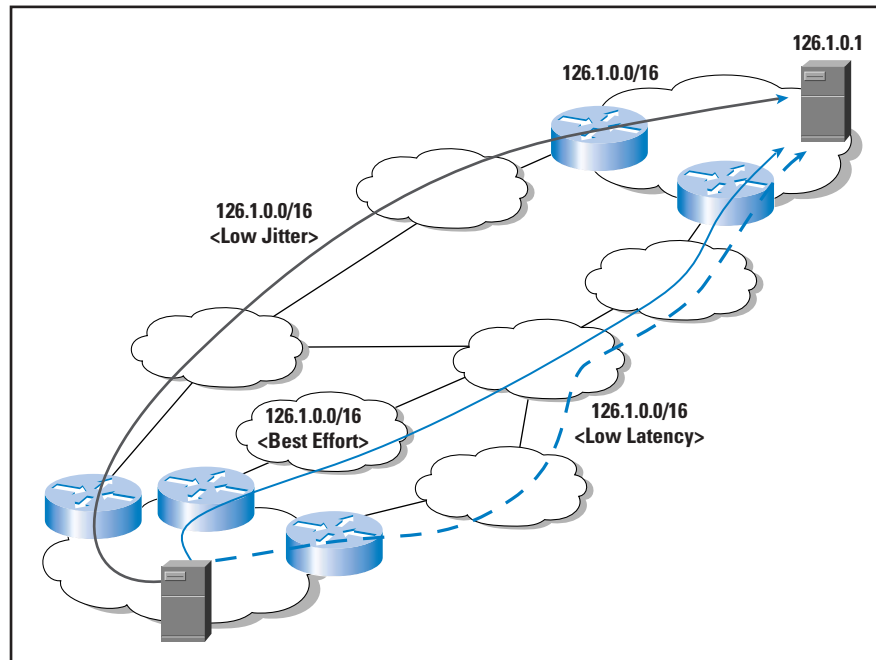
At an abstract level, this greater volume of routing information is needed in order to address the dual role of the routing system as both an inter-domain connectivity maintenance protocol and as a traffic-engineering tool.

Inter-Domain Quality of Service

Quality of Service (QoS) is a term that encompasses a wide variety of mechanisms. In the case of routing, the term is used to describe the process of modifying the normal routing response of associating a single forwarding action with a destination address prefix in such a way that there may be numerous forwarding decisions for a particular address prefix. Each forwarding decision is associated with a particular service response, so that a “best-effort” path to a particular destination address may differ from a “low-latency” path, which in turn may differ from a “high-bandwidth” path, and so on.

As with inter-domain traffic engineering, this requirement is one which would be expected to place greater volumes of information into the routing domain. At an abstract level this requirement can be seen as the association of a service quality attribute with an address prefix, and passing the paired entity into the routing domain as a single routing object. The inference is that multiple quality attributes associated with a path to a particular prefix would require the routing system to independently manipulate multiple route objects, because it would be reasonable to anticipate that the routing system would select different paths to reach the same address prefix if different QoS service attributes were used as a path qualifier (Figure 3).

Figure 3: Inter-Domain
Routing with QoS



Approaches to Inter-Domain Routing

Let's now take this set of requirements and attempt to match them to various approaches to routing protocols.

Routing is a distributed computation wherein each element of the computation set must reach an outcome that is consistent with all other computations undertaken by other members of the set. There are two major approaches to this form of distributed computation, namely *serial* or *parallel* computation. Serial computation involves each element of the set undertaking a local computation and then passing the outcomes of this computation to its adjacent elements. This approach is used in various forms of distance-vector routing protocols where each routing node computes a local set of selected paths, and then propagates the set of reachable prefixes and the associated path metric to its neighbors. Parallel computation involves rapid flooding of the current state of connectivity within the set to all elements, and all set elements simultaneously compute forwarding decisions using the same base connectivity data. This approach is used in various forms of link-state routing protocols, where the protocol uses a flooding technique to rapidly propagate updated link-status information and then relies on each routing node to perform a local path selection computation for each reachable address prefix. Is one of these approaches substantially better suited than the other to the inter-domain routing environment?

Open or Closed Routing Policies

One of the key issues behind consideration of this topic is that of the role of *local policy*. Using a distance-vector protocol, a routing domain gathers selected path information from its neighbors, applies local policy to this information, and then distributes this updated information in the form of selected paths to its neighbor domains.

In this model the nature of the local policy applied to the routing information is not necessarily visible to the domain neighbors, and the process of converting received route advertisements into advertised route advertisements uses a local policy process whose policy rules are not visible externally. This scenario can be described as *policy opaque*. The side effect of such an environment is that a third party cannot remotely compute which routes a network may accept and which may be readvertised to each neighbor.

In link-state protocols, a routing domain effectively broadcasts its local domain adjacencies, and the policies it has with respect to these adjacencies, to all nodes within the link-state domain. Every node can perform an identical computation upon this set of adjacencies and associated policies in order to compute the local inter-domain forwarding table. The essential attribute of this environment is that the routing node has to announce its routing policies in order to allow a remote node to compute which routes will be accepted from which neighbor, and which routes will be advertised to each neighbor and what, if any, attributes are placed on the advertisement. Within an interior routing domain the local policies are in effect metrics of each link, and these policies can be announced within the routing domain without any consequent impact.

In the exterior routing domain it is not the case that interconnection policies between networks are always fully transparent. Various permutations of supplier/customer relationships and peering relationships have associated policy qualifications that are not publicly announced for business competitive reasons. The current diversity of interconnection arrangements appears to be predicated on policy opacity, and to mandate a change to a model of open interconnection policies may be contrary to operational business imperatives. An inter-domain routing tool should be able to support models of interconnection where the policy associated with the interconnection is not visible to any third party. If the architectural choice is a constrained one between distance vector and link state, then this consideration would appear to favor the continued use of a distance-vector approach to inter-domain routing. This choice, in turn, has implications on the convergence properties and stability of the inter-domain routing environment. If there is a broader spectrum of choice, the considerations of policy opacity would still apply.

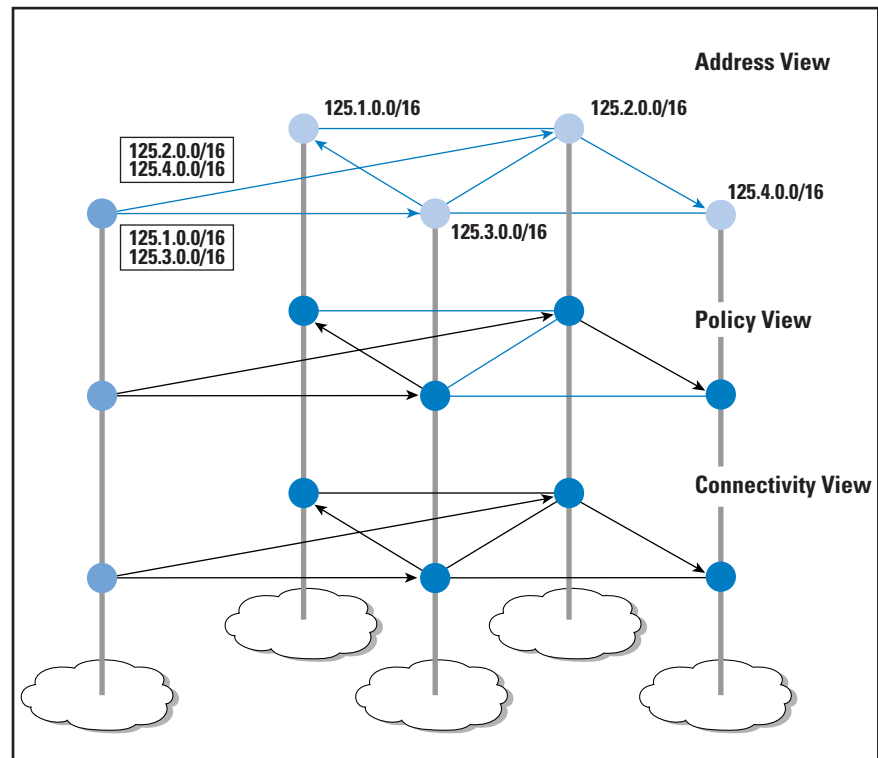
Separation of Functions

The inter-domain routing function undertakes many roles simultaneously. First, it maintains the current view of inter-domain connectivity. Any changes in the adjacency of a domain are reflected in a distributed update computation that determines if the adjacency change implies a change in path selection and in address reachability. Secondly, it maintains the set of currently reachable address prefixes. And finally, the protocol binds the first two functions together by associating each prefix with a path through the inter-domain space.

This association uses a policy framework to allow each domain to select a path that optimizes local policy constraints within the bounds of existing constraints applied by other domains. This policy may be related to traffic-engineering objectives, QoS requirements, local cost optimization, or related operational or business objectives.

An alternative approach to inter-domain routing is to separate the functions of connectivity maintenance, address reachability, and policy negotiation. As an example of this approach, a connectivity protocol can be used to identify all viable paths between a source and a destination domain. A policy negotiation protocol can be used to ensure that there are a consistent sequence of per-domain forwarding decisions that will pass traffic from the source domain to the destination domain. An address reachability protocol can be used to associate a collection of address prefixes with each destination domains. This framework is illustrated in Figure 4.

Figure 4: A Multi-Tiered Approach to Inter-Domain Routing



Address Prefixes and Autonomous System Numbers

One observation about the current inter-domain routing system is that it uses a view of the network based on computing the optimal path to each address prefix. This view is translated into an inter-domain routing protocol that uses the address prefix as the basic protocol element and attaches various attributes to each address prefix as they are passed through the network

As of late 2001, the routing system had some 100,000 distinct address prefixes and 11,500 origin domains. This implies that each origin domain is responsible for an average of 8 to 9 address prefixes. If each domain advertised its prefixes with a consistent policy, then each address prefix would be advertised with identical attributes. If the routing protocol were to be inverted such that the routing domain identifier, or *Autonomous System* number, were the basic routing object and the set of prefixes and associated common set of route attributes were attributes of the Autonomous System object, then the number of routing objects would be reduced by the same factor of between 8 and 9.

The motivation in this form of approach is that seeking clear hierarchical structure in the address space as deployed is no longer feasible, and that no further scaling advantage can be obtained by various forms of address aggregation within the routing system. This approach replaces this address-based hierarchy with a two-level hierarchy of routing domains. Within a routing domain, routing is undertaken using the address prefix. Between routing domains, routing is undertaken using domain identifiers and associated sets of domain attributes.

Although this approach appears to offer some advantage in creating a routing domain, one-tenth of the size of the address prefix-based routing domain, it is interesting to note that since late 1996 the average number of address prefixes per Autonomous System has fallen from 25 to the current value of 9. In other words, the number of distinct routing domains is growing at a faster rate than the number of routed address prefixes. While the adoption of a domain-based routing protocol offers some short-term advantages in scaling, the longer-term prospects are not so attractive, given these relative growth rates.

Routing Hierarchies of Information

The scaling properties of an inter-domain routing protocol are related on the ability of the protocol to remove certain specific items of information from the routing domain at the point where it ceases to have any differentiating impact. For example, it is important for a routing protocol to carry information that a particular domain has multiple adjacencies and that there are a number of policies associated with each adjacency, and propagate this information to all local domains. At a suitably distant point in the network, the forwarding decision remains the same regardless of the set of local adjacencies, and propagation of the detail of the local environment to points where the information ceases to have any distinguishing outcome is unproductive.

From this perspective, scaling the routing system is not a case of determining what information can be added into the routing domain, but instead it's a case of determining how much information can be removed from the routing domain, and how quickly.

One way of removing information is through the use of *hierarchies*. Within a hierarchical structure, a set of objects with similar properties are aggregated into a single object with a set of common properties. One way to perform such aggregation is by increasing the amount of information contained in each aggregate route object. For example, if single route objects are to be used that encompass a set of address prefixes and a collection of Autonomous Systems, then it would be necessary to define additional attributes within the route object to further qualify the policies associated with the object in terms of specific prefixes, specific Autonomous Systems, and specific policy semantics that may be considered as policy exceptions to the overall aggregate. This approach would allow aggregation of routing information to occur at any point in the network, allowing the aggregator to create a compound object with a common set of attributes, and a set of additional attributes that apply to a particular subset of the aggregate.

Another approach to using hierarchies to reduce the number of route objects is to reduce the scope of advertisement of each routing object, allowing the object to be removed and proxy aggregated into some larger object when the logical scope of the object is reached. This approach would entail the addition of route attributes that could be used to define the circumstances where a specific route object would be subsumed by an aggregate route object without impacting the policy objectives associated with the original set of advertisements. This approach places control of aggregation with the route object originator, allowing the originator to specify the extent to which a specific route object should be propagated before being subsumed into an aggregate object.

It is not entirely clear that the approach of exploiting hierarchies in an address space is the most appropriate response to scaling pressures. Viewed from a more general perspective, scaling of the routing system requires the systematic removal of information from the routing domain. The way this is achieved is by attempting to align the structure of deployment with some structural property of the syntax of the protocol elements that are being used as routing objects. Information can then be eliminated through systematic aggregation of the routing objects at locations within the routing space that correspond to those points in the topology of the network where topology aggregation is occurring. The maintenance of this tight coupling of the structure of the deployed network to the structure of the identifier space is the highest cost of this approach. Alterations to the topology of the network through the relocation or reconfiguration of networks requires renumbering of the protocol element if hierarchical aggregation is to be maintained. If the address space is the basis of routing, as at present, then this becomes a large-scale exercise of renumbering networks that in turn implies an often prohibitively disruptive and expensive exercise of renumbering collections of host systems and associated services.

One view of this is that the connectivity properties of the Internet are already sufficiently meshed that there is no readily identifiable hierarchical structure, and that this trend is becoming more pronounced, not less. In that case, the most appropriate course of action may be to reexamine the routing domain and select some other attribute as the basis of the routing computation that does not have the same population, complexity, and growth characteristics as address prefixes, and base the routing computation on this attribute. One such alternative approach is to consider Autonomous System numbers as routing “atoms” where the routing system converges to select an Autonomous System path to a destination Autonomous System, and then uses this information to add the associated set of prefixes originated by this Autonomous System, and next-hop forwarding decision to reach this Autonomous System into the local forwarding table.

Extend or Replace BGP

A final consideration is to consider whether these requirements can best be met by an approach of a set of upward-compatible extensions to BGP, or by a replacement to BGP.

The rationale for extending BGP would be to increase the number of commonly supported transitive route attributes, and, potentially, allow a richer syntax for attribute definition which in turn would allow the protocol to use a richer set of semantic definitions in order to express more complex routing policies.

This direction may sound like a step backward, in that it proposes an increase in the complexity of the route objects carried by the protocol and potentially increases the amount of local processing capability required to generate and receive routing updates. However, this can be offset by potential benefits that are realizable through the greater expressive capability for the policy attributes associated with route objects. It can allow a route originator an ability to specify the scope of propagation of the route object, rather than assuming that propagation will be global. The attributes can also describe intended service outcomes in terms of policy and traffic engineering. It may also be necessary to allow BGP sessions to negotiate additional functionality intended to improve the convergence behavior of the protocol. Whether such changes can produce a scalable and useful outcome in terms of inter-domain routing remains, at this stage, an open question.

An alternative approach is that of a replacement protocol. Use of a parallel-processing approach to the distributed computation of routing, such as that used in the link-state protocols, can offer the benefits of faster convergence times and avoidance of unstable transient routing states. On the other hand, link-state protocols present issues relating to policy opaqueness, as described above. Another major issue with such an approach is the need to address the efficiency of inter-domain link-state flooding.

The inter-domain space would need some further levels of imposed structure similar to intra-domain areas in order to ensure that individual link updates are rapidly propagated across the relevant subset of the network. The use of such an area structure may well imply the need for an additional set of operator relationships, such as mutual transit. Such inter-domain relationships may prove challenging to adapt to existing operator practices.

Another approach could be based on the adoption of a multi-layer approach of separate protocols for separate functions, as described above. A base inter-domain connectivity protocol could potentially be based on a variant of a link-state protocol, using the rapid convergence properties of such protocols to maintain a coherent view of the current state of connectivity within the network. The overlay of a policy protocol would be intended as a signaling mechanism to allow each domain to make local forwarding decisions that are consistent with those adopted by adjacent domains, thereby maintaining a collection of coherent inter-domain paths from source to destination. Traffic engineering can also be envisaged as an overlay mechanism, allowing a source to make a forwarding decision that selects a path to the destination where the characteristics of the path optimize the desired service outcomes.

Directions for Further Activity

Although short-term actions based on providing various incentives for network operators to remove redundant or inefficiently grouped entries from the BGP routing table may exist, such actions are short-term palliative measures, and will not provide long-term answers to the need for a scalable inter-domain routing protocol. One approach to the longer-term requirements may be to preserve many of the attributes of the current BGP protocol, while refining other aspects of the protocol to improve its scaling and convergence properties. A minimal set of alterations could retain the Autonomous System concept to allow for administrative boundaries of information summarization, as well as retaining the approach of associating each prefix advertisement with an originating Autonomous System. The concept of policy opaqueness would also be retained in such an approach, implying that each Autonomous System accepts a set of route advertisements, applies local policy constraints, and readvertises those advertisements permitted by the local policy constraints. It could be feasible to consider alterations to the distance-vector path-selection algorithm, particularly as it relates to intermediate states during processing of a route withdrawal. It is also feasible to consider the use of compound route attributes, allowing a route object to include an aggregate route, and numerous specifics of the aggregate route, and attach attributes that may apply to the aggregate or a specific address prefix. Such route attributes could be used to support multi-homing and inter-domain traffic-engineering mechanisms. The overall intent of this approach is to address the major requirements in the inter-domain routing space without using an increasing set of globally propagated specific route objects.

Another approach is to consider the feasibility of decoupling the requirements of inter-domain connectivity management with the applications of policy constraints and the issues of sender- and receiver-managed traffic-engineering requirements. Such an approach may use a link-state protocol as a means of maintaining a consistent view of the topology of inter-domain network, and then use some form of overlay protocol to negotiate policy requirements of each Autonomous System, and use a further overlay to support inter-domain traffic-engineering requirements. The underlying assumption of such an approach is that if the functional role of inter-domain routing is divided into distinct components, each component will have superior scaling and convergence properties which in turn will result in superior properties for the entire routing system. Obviously, this assumption requires some testing.

Research topics with potential longer-term application include the approach of drawing a distinction between the identity of a network, its location relative to other networks, and maintenance of a feasible path set between a source and destination network that satisfies various policy and traffic-engineering constraints. Again the intent of such an approach would be to divide the current routing function into numerous distinct scalable components rather than using a single monolithic routing protocol.

Further Reading

- [0] Huston, G., "Analyzing the Internet BGP Routing Table," *The Internet Protocol Journal*, Vol. 4, No. 1, March 2001.
www.cisco.com/warp/public/759/ipj_4-1/ipj_4-1_bgp.html
- [1] Huitema, C., *Routing in the Internet, 2nd Edition*, ISBN 0130226475, Prentice Hall, January 2000. *A good introduction to the general topic of IP routing.*
- [2] Rekhter, Y., and Li T., "A Border Gateway Protocol 4 (BGP-4)," RFC 1771, March 1995. *The base specification of BGP 4. This document is currently being updated by the IETF. The state of this work in progress as of November 2001 is documented as an Internet Draft,*
[draft-ietf-idr-bgp4-15.txt](#)
- [3] Elwyn Davies et al., "Future Domain Routing Requirements," work in progress, July 2001. *This work is currently documented as an Internet Draft, **[draft-davies-fdr-reqs-01.txt](#)**. It contains a review of an earlier effort in enumerating routing requirements ("Goals and Functional Requirements for Inter-Autonomous System Routing," RFC 1126, October 1989), as well as a commentary on a proposed set of current routing requirements.*

GEOFF HUSTON holds a B.Sc. and M.Sc. from the Australian National University. He has been closely involved with the development of the Internet for the past decade, particularly within Australia, where he was responsible for the initial build of the Internet within the Australian academic and research sector. He is currently the Chief Scientist in the Internet area for Telstra, a member of the Internet Architecture Board, and is the Secretary of the APNIC Executive Committee. He is author of *The ISP Survival Guide*, ISBN 0-471-31499-4, *Internet Performance Survival Guide: QoS Strategies for Multiservice Networks*, ISBN 0471-378089, and coauthor of *Quality of Service: Delivering QoS on the Internet and in Corporate Networks*, ISBN 0-471-24358-2, a collaboration with Paul Ferguson. All three books are published by John Wiley & Sons. E-mail: **gih@telstra.net**

Development of the Regional Internet Registry System

by Daniel Karrenberg, RIPE-NCC; Gerard Ross, APNIC; Paul Wilson, APNIC; Leslie Nobile, ARIN

The current system of managing Internet address space involves *Regional Internet Registries* (RIRs), which together share a global responsibility delegated to them by the *Internet Assigned Numbers Authority* (IANA). This regime is now well established, but it has evolved over ten years from a much simpler, centralized system. Internet number spaces were originally managed by a single individual “authority,” namely the late Jon Postel, co-inventor of some of the most important technical features of today’s Internet.

It is important to understand that the evolution of the RIR system was not simply the result of Internet growth and the natural need to refine and decentralize a growing administrative task. On the contrary, it arose from, and closely tracked, the technical evolution of the Internet Protocol, in particular the development of today’s IP addressing and routing architecture.

In a relatively short time, the Regional Internet Registry system has evolved into a stable, robust environment for Internet address management. It is maintained today through self-regulatory practices that are well established elsewhere in the Internet and other industries, and it maintains its legitimacy and relevance by firmly adhering to open, transparent, participatory decision-making processes.

Before the RIRs:

IP Address Architecture

An important feature of the Internet Protocol (IP) is the ability to transparently use a wide variety of underlying network architectures to transport IP packets. This is achieved by encapsulating IP packets in whatever packet or frame structure the underlying network uses. Routers connecting different networks forward IP traffic by decapsulating incoming IP packets and then re-encapsulating them as appropriate for the next network to carry them.

To achieve this task with full transparency, the IP needed an addressing structure, which developed as a two-level hierarchy in both addressing and routing. One part of the address, the *network* part, identifies the particular network a host is connected to, while the other part, the *local* part, identifies the particular end system on that network.

Internet routing, then, has to deal only with the network part of the address, routing the packet to a router directly connected to the destination network. The local part is not used at all in Internet routing itself; rather it is used to determine the intended address within the addressing structure of the destination network.

The method by which the local part of an IP address is translated to a local network address depends on the architecture of the destination network—static tables, simple conversions, or special-purpose protocols are used as appropriate.

The original Internet addresses comprised 32 bits, the first 8 bits providing the network part and the remaining 24 bits the local part. These addresses were used for many years. However, in June 1978, in Internet Engineering Note (IEN) 46 “A proposal for addressing and routing in the internet,” Clark and Cohen observed:

“The current internet header has space to name 256 networks. The assumption, at least for the time being, is that any network entering the internet will be assigned one of these numbers. While it is not likely that a great number of large nets, such as the ARPANET, will join the internet, the trend toward local area networking suggests that a very large number of small networks can be expected in the internet in the not too distant future. We should thus begin to prepare for the day when there are more than 256 networks participating in the internet.”

Classful Addressing

As predicted, it soon became necessary to adapt the address architecture to allow more networks to be connected. By the time the Internet Protocol itself was comprehensively specified (in RFC 790, published in 1981, edited by Jon Postel), the IP address could be segmented in numerous ways to provide three classes of network address.

In Class A, the high-order bit is zero, the next 7 bits are the network, and the last 24 bits are the local address. In Class B, the high-order 2 bits are one-zero, the next 14 bits are the network, and the last 16 bits are the local address. In Class C, the high-order 3 bits are one-one-zero, the next 21 bits are the network, and the last 8 bits are the local address.

This so-called “classful” architecture served the Internet for the next 12 years, during which time it grew from a small U.S.-based research network to a global academic network showing the first signs of commercial development.

Early Registration Models

In the 1980s, the American *National Science Foundation’s* (NSF’s) high-speed network, NSFNET, was connected to the ARPANET, a U.S. *Defense Advanced Research Projects Agency* (ARPA, now DARPA) wide-area network, which essentially formed the infrastructure that we now know as the Internet.

From these early days of the Internet, the task of assigning addresses was a necessary administrative duty, to ensure simply that no two networks would attempt to use the same network address in the Internet.

At first, the elementary task of maintaining a list of assigned network addresses was carried out voluntarily by Jon Postel, using (according to legend) a paper notebook.

As the Internet grew, and particularly as classful addressing was established, the administrative task grew accordingly. The IANA was established, and within it the Internet Registry (IR). But as the task of the IR outgrew Postel's notebook, it was passed to SRI International in Menlo Park, California, under a NSF contract, and was called the *Defense Data Network (DDN) Network Information Center (NIC)*.

During this time, under the classful address architecture, networks were allocated liberally and to any organization that fulfilled the simple request requirements. However, with the accelerating growth of the Internet during the late 1980s, two problems loomed: the rapid depletion of address space, due to the crude classful divisions; and the uncontrolled growth of the Internet routing table, due to unaggregated routing information.

Conservation vs. Aggregation

The problems of “three sizes fit all” highlight the basic dilemma of address space assignment: conservation versus aggregation. On the one hand, one wants to conserve the address space by assigning as little as possible; on the other hand, one wants to ease routing-table pressures by aggregating as many addresses as possible in one routing-table entry.

This can be illustrated by looking at a typical networking setup of the time. Within organizations having a single Internet connection, buildings, departments, or campuses would have their own local networks. Often the use of multiple networks was dictated by distance limitations inherent in the emerging local-area networking technologies, such as Ethernet.

These networks typically had to accommodate more than the 254 hosts addressable by a Class C address, but would rarely exceed 1000 hosts. Using pure classful addressing, one could either subdivide networks artificially to remain below the 254 host limit, or use a Class B address for each local network, possibly wasting more than 60,000 addresses in each. Whereas the latter solution is obviously wasteful in terms of address space, the former is obviously cumbersome. Less obviously, the former also puts an additional burden on the Internet routing system, because each of these networks would require a separate route propagated throughout the whole Internet.

This basic dilemma persists to this day. Assigning address space generously tends to reduce the routing-table size, but wastes address space. Assigning conservatively will waste less, but cause more stress for the routing system.

Subnetting

In order to address some of the problems of classful addressing, the technique of *subnetting* was invented. Described in RFC 791 in 1984, subnetting provided another level of addressing hierarchy by inserting a *subnet* part into the IP address between the network and local parts. Global routing remained the same using the *network* part of the address (Class A, B, or C) until traffic reached a router on the network identified by the network part of the address. This router, configured for subnetting, would interpret a statically configured number of bits from the local part of the address (the subnet part) to route the packet further among a set of similarly configured routers. When the packet reached a router connected to the destination subnet, the remaining bits of the local part would be used to determine the local address of the destination as usual. So, in the previous example, the organization could have used a Class B address with 6-bit subnetting, a setup that would allow for 62 networks of 1022 hosts each.

Subnetting nicely solved the routing-table problem, because now only one global routing-table entry was needed for the organization. It also helped address space conservation somewhat because it provided an obvious alternative to using many sparsely populated Class B networks.

Because the boundary between the subnet part and the local part of an address could not be determined from the address itself, this local knowledge needed to be configured into the routers. At first this was done by static configuration. Later, interior routing protocols carried that information. Refer to RFC 791 for numerous historically interesting case studies.

Supernetting

Within seven years, however, it was becoming clear that subnetting was no longer sufficient to keep up with Internet growth. RFC 1338 stated the problem:

“As the Internet has evolved and grown ... in recent years, it has become painfully evident that it is soon to face several serious scaling problems. These include:

1. Exhaustion of the Class-B network address space. One fundamental cause of this problem is the lack of a network class of a size that is appropriate for a midsized organization; Class C, with a maximum of 254 host addresses, is too small while Class B, which allows up to 65534 addresses, is too large to be widely allocated.
2. Growth of routing tables in Internet routers beyond the ability of current software (and people) to effectively manage.
3. Eventual exhaustion of the 32-bit IP address space.

It has become clear that the first two of these problems are likely to become critical within the next one to three years.”

The solution proposed was to extend the subnetting technique beyond the local organization, into the Internet itself. In other words, RFC 1338 proposed abolishing classful addressing, and replacing it with *supernetting*. The proposal was summarized as follows:

“The proposed solution is to hierarchically allocate future IP address assignment, by delegating control of segments of the IP address space to the various network service providers.”

CIDR

In 1993, the supernetting technique was published as a standards track RFC under the name *Classless Inter-Domain Routing* (CIDR), by which it is known and used today. Two main ingredients were necessary to make CIDR work: routing system changes and new address allocation and assignment procedures.

Under CIDR, routers could no longer determine the network part of an address from the address itself. This information now needed to be conveyed by Internet routing protocols. Fortunately, there was only one such protocol in widespread use at the time, and it was quickly extended by the major router vendor of the time. According to legend, the necessary extensions of the *Border Gateway Protocol* (BGP)-3 to BGP-4 were designed on a napkin, with all implementors of significant routing software present. The changes were implemented in a matter of days, but only much later described by the Internet standards track RFC 1654.

CIDR also required that forwarding decisions of routers be changed slightly. The network part of an address, now more generally called the *prefix*, can be of any length. This means that a router can have multiple valid routes covering a specific 32-bit destination address. Routers need to use the most specific of these routes—the *longest prefix*—when forwarding packets.

In addition to technical changes, the success of CIDR also relied on the development of administrative procedures to allocate and assign address space in such a way that routes could be aggregated as much as possible. Because the Internet was evolving toward the current state of arbitrarily interconnected networks of *Internet Service Providers* (ISPs), it was obvious that ISPs should play a role in address space distribution. In the new technique, ISPs would now, as much as possible, assign address space to their customers in contiguous blocks, which could be aggregated into single routes to the rest of the Internet.

Emergence of the RIRs:

Internationalization

While the engineering-driven need for topological address space assignment was becoming clear, there was also an emerging recognition that the administrative mechanisms of address space distribution needed further development. A central system just would not scale for numerous reasons, including:

- Sheer volume
- Distance from the address space consumers
- Lack of an appropriate global funding structure
- Lack of local community support

The need to change administrative procedures was formally recognized by August 1990, when the Internet Activities Board published a message it had sent to the U.S. Federal Networking Council, stating “it is timely to consider further delegation of assignment and registration authority on an international basis” (RFC 1174).

The increasing cultural diversity of the Internet also posed administrative challenges for the central IR. In October 1992, the *Internet Engineering Task Force* (IETF) published RFC 1366, which described the “growth of the Internet and its increasing globalization” and set out the basis for an evolution of the registry process, based on a regionally distributed registry model. This document stressed the need for a single registry to exist in each geographical region of the world (which would be of “continental dimensions”). Registries would be “unbiased and widely recognized by network providers and subscribers” within their region. Each registry would be charged with allocating remaining address space in a manner “compatible with potential address aggregation techniques” (or CIDR).

RIPE NCC

While in the United States the Government continued to support and fund registry functions, this was not the case in other parts of the world. In Europe, IP network operators cooperating in *Réseaux IP Européens* (RIPE) realized the need for professional coordination and registration functions. Establishment of the *RIPE Network Coordination Centre* (NCC) was proposed in the same month that RFC 1174 was published. The RIPE NCC was to “function as a ‘Delegated Registry’ for IP numbers in Europe, as anticipated and defined in RFC 1174” (RIPE-19).

Although consensus among IP network operators was quickly established, it took almost two years of organizing and fund-raising before the first RIR was fully operational in May 1992. The RIPE NCC was organized as a highly independent part of RARE, the organization of European research networks. It was to be funded by contributions from those networks, as well as a small number of emerging commercial networks. The RIPE NCC published its first regional address distribution policy in July 1992 (RIPE-65).

During the following months, European regional policies were refined and, for the first time, global guidelines were published as RFCs (RFC 1366, RFC 1466).

The RIPE NCC is presently organized as a membership association, performing the essential coordination and administration activities required by the RIPE community. Located in Amsterdam, Netherlands, the RIPE NCC service region incorporates 109 countries covering Europe, the Middle East, Central Asia, and African countries located north of the equator. The RIPE NCC currently consists of more than 2700 members. At the time of publication, RIPE NCC is performing the secretariat function for the *Address Supporting Organization* (ASO) of The *Internet Corporation for Assigned Names and Numbers* (ICANN). More information about RIPE NCC is available at <http://www.ripe.net>

APNIC

Asia Pacific Network Information Centre (APNIC), the second RIR, was established in Tokyo in 1993, as a pilot project of APCCIRN (Asia Pacific Coordination Committee for Intercontinental Research Networks, now *Asia Pacific Networking Group* [APNG]).

The project was an intended as a trial model for servicing the Internet addressing needs of national *Network Information Centres* (NICs) and other networks throughout the region.

After a successful ten-month trial period, APNIC was established as a permanent organization to serve the Asia Pacific region (which includes 62 economies from Central and South Asia to the Islands of Oceania and the Western Pacific).

Originally, APNIC relied on the support of networking organizations and national NICs. However, in 1996, APNIC implemented a tiered membership structure.

APNIC relocated to Brisbane, Australia, in mid-1998. It currently services approximately 700 member organizations, across 39 economies of the region. Within the APNIC membership, there are also five *National Internet Registries* (NIRs), in Japan, China, Taiwan, Korea, and Indonesia. The NIRs perform analogous functions to APNIC at a national level and together represent the interests of more than 500 additional organizations.

In 2000, APNIC hosted the secretariat functions of the ASO in its inaugural year. More information about APNIC is available at:

<http://www.apnic.net>

ARIN

In 1991, the contract to perform the IR function was awarded to Network Solutions, Inc. in Herndon, Virginia. This included the transition of services including IP address registration, domain name registration and support, *Autonomous System Number* (AS) registration, user registration, online information services, help-desk operations, and RFC and Internet-Draft archive and distribution services (RFC 1261).

With explosive Internet growth in the early 1990s, the U.S. Government and the NSF decided that network support for the commercial Internet should be separated from the U.S. Department of Defense. The NSF originated a project named InterNIC under a cooperative agreement with *Network Solutions, Inc.* (NSI) in 1993 to provide registration and allocation of domain names and IP address numbers for Internet users.

Over time, after lengthy consultation with the IANA, the IETF, RIPE NCC, APNIC, the NSF, and the *Federal Networking Council* (FNC), a further consensus was reached in the general Internet community to separate the management of domain names from the management of IP numbers. This consensus was based on the recognition that the stability of the Internet relies on the careful management of IP address space.

Following the examples of RIPE NCC and APNIC, it was recommended that management of IP address space then administered by the InterNIC should be under the control of, and administered by, those that use it, including ISPs, end-user organizations, corporate entities, universities, and individuals.

As a result, ARIN (*American Registry for Internet Numbers*) was established in December 1997, as an independent, nonprofit corporation, with a membership structure open to all interested entities or individuals.

ARIN is located in Chantilly, Virginia, United States. Its service region incorporates 70 countries, covering North America, South America, the Caribbean, and African countries located south of the equator. ARIN currently consists of more than 1500 members. Within the ARIN region, there are two national delegated registries, located in Mexico and Brazil.

Until now, ARIN has carried the responsibility for maintaining registration of resources allocated before the inception of the RIRs. However, a major project is now under way to transfer these legacy records to the relevant RIRs. More information about ARIN is available at:

<http://www.arin.net>

Emerging RIRs

The existing RIRs currently serve countries outside their core regions to provide global coverage; however, new RIRs are expected to emerge, necessitating changes to the existing service regions. Because the regions are defined on continental dimensions, the number of new RIRs will be low.

Currently, two groups have made significant progress in seeking to establish new RIRs. *AfriNIC* (for the Africa region) and *LACNIC* (for Latin America and the Caribbean) have each conducted public meetings, published documentation, and participated in the activities of the

existing RIRs. In recognition of the regional support they have so far obtained, each organization has been granted observer status at ICANN ASO meetings. The existing RIRs have also sought to provide as much assistance and support as possible to these emerging organizations.

More information about AfriNIC is available at;
<http://www.afrinic.org/>

More information about LACNIC is available at:
<http://lacnic.org/>

The RIR System:

Goals of the RIRs

RFC 2050, published in November 1996, represented a collaboration of the global Internet addressing community to describe a set of goals and guidelines for the RIRs. Although IANA was to retain ultimate responsibility for the entire address pool, RFC 2050 recognizes that RIRs operate under the consensus of their respective regional Internet community. This document, along with a history of RIR coordination, has helped to form the basis for a set of consistent global policies.

The three primary goals of the RIR system follow:

- *Conservation*: to ensure efficient use of a finite resource and to avoid service instabilities due to market distortions (such as stockpiling or other forms of manipulation);
- *Aggregation (routability)*: to assist in maintenance of Internet routing tables at a manageable size, by supporting CIDR techniques to ensure continued operational stability of the Internet;
- *Registration*: to provide a public registry documenting address space allocations and assignments, necessary to ensure uniqueness and provide information for Internet troubleshooting at all levels.

The Open Policy Framework

It was always recognized that these goals would often be in conflict with each other and with the interests of individuals and organizations. It was also recognized that legitimate regional interests could justify varying approaches in balancing these conflicts. Therefore, within the global framework, each regional community has always developed its own specific policies and procedures.

However, whereas the specific approaches may differ across the RIRs, all operate on a basic principle of open, transparent, consensus-based decision-making, following self-regulatory practices that exist elsewhere in the Internet and other industries. Furthermore, the RIRs all maintain not-for-profit cost-recovery systems and organizational structures that seek to be inclusive of all interested stakeholders.

The activities and services of each of the RIRs are defined, performed, discussed, and evaluated in open forums, whose participants are ultimately responsible for decision-making.

To facilitate broad participation, open policy meetings are hosted by RIRs regularly in each of the regions. Ongoing discussions are carried out on the public mailing lists of each RIR, which are open to both the RIR constituents and the broader community. The RIRs also participate actively in other Internet conferences and organizations and, importantly, each RIR has a strong tradition of participating in the public activities of the others.

A current example of the coordinated efforts of the RIRs is the Provisional IPv6 Assignment and Allocation Policy Document, a joint effort of the RIRs with the assistance of the IETF, The *Internet Architecture Board* (IAB), and the *Internet Engineering Steering Group* (IESG) to describe the allocation and assignment policies for the first release of IPv6 address numbers.

Also, the RIRs recently published the RIR Comparative Policy Overview, which is available at: <http://www.ripe.net/ripencc/mem-services/registration/rir-comp-matrix-rev.html>

These documents help illustrate that the well-established combination of bottom-up decision-making and global cooperation of the RIRs has created a stable, robust environment for Internet address management.

RIR Functions

The primary function of each RIR is to ensure the fair distribution and responsible management of IP addresses and the related numeric resources that are required for the stable and reliable operation of the Internet. In particular, the resources allocated, assigned, and registered by RIRs are Internet address numbers (IPv4 and IPv6) and AS numbers. RIRs are also responsible for maintaining the reverse delegation registrations of the parent blocks within their respective ranges.

Complementing their registry function, the RIRs have an important role in educating and informing their communities. The activities carried out by the individual RIRs vary, but include open policy meetings, training courses, seminars, outreach activities, statistical reporting, and research.

Additionally, a crucial role for the RIRs is to represent the interests of their communities by participating in global forums and providing support to other organizations involved in Internet addressing issues.

RIRs and The Global Internet Community:

Formation of ICANN and the ASO

The global Internet governance landscape began to undergo radical changes in mid-1998, with the publication of a U.S. Government white paper outlining the formation of a “not-for-profit corporation formed by private sector Internet stakeholders to administer policy for the Internet name and address system.” ICANN was formed later that year.

At the heart of the ICANN structure are “supporting organizations” that are formed to “assist, review and develop recommendations on Internet policy and structure” within specialized areas. In October 1999, the existing RIRs and ICANN jointly signed a *Memorandum of Understanding* (MoU) to establish the principles for forming and operating the *Address Supporting Organization* (ASO). It is intended that new RIRs will sign the MoU as they emerge.

Under the ASO MoU, the policy forums within each of the RIR regions continue to be responsible for development of regional IP address policy. In addition, each signatory RIR is responsible for electing three members to the ICANN *Address Council*.

The purpose of the Address Council, as described in the MoU, is to review and develop recommendations on issues related to IP address space, using the open processes that exist in the three regions; and to advise the ICANN Board on these matters. In addition, the Address Council is responsible for the appointment of three ICANN Directors to the ICANN Board.

RIR–ASO Coordination

Since the formation of the ASO, the RIRs have played an integral part in facilitating its activities. By joint agreement, the RIRs will share the ASO secretariat duties, including the hosting of the ASO Web site, on a revolving basis. APNIC provided these services in the ASO’s first year of operation, and RIPE NCC is currently performing this role.

The ASO Address Council holds monthly telephone conferences, which are attended by representatives of the RIRs (and emerging RIRs on a listener basis). In accordance with the MoU, the ASO also holds regular open meetings in conjunction with the open policy meetings of the RIRs.

RIRs and Industry Development

As noted previously, the RIRs maintain high levels of participation in the conferences and activities of other organizations. Similarly, they invite the participation of interested parties in their own activities.

The RIRs are active in many areas of new technology implementation (such as *General Packet Radio Service* [GPRS] and *Universal Telecommunications System* [UMTS] mobile telephony, IPv6, and cable and *Digital Subscriber Line* [xDSL]-based Internet services).

The established regional processes have proved both flexible and open enough to incorporate such new developments into policy formation. Industry representatives frequently join policy discussions, present at plenary sessions, and participate in working groups.

The RIRs pursue relationships with industry bodies, particularly those with representative and developmental functions, to facilitate industry convergence on open standards and policy processes.

Many diverse parties have legitimate interests in the allocation and registration of IP addresses, and the RIRs remain committed to participating with these parties to achieve a consensus among the Internet community on IP address allocation issues.

The Future of RIRs

In Internet time it can be easy to forget that eight years is actually not long. Since it was first proposed in 1990, the RIR system has evolved rapidly, enjoyed strong community support, and has been relatively free of the political wrangling that has characterized the registration systems of other Internet resources. Without doubt, this position is largely due to the early determination to provide accessible, open forums for the interested stakeholders in the various regions.

New technologies, such as GPRS, broadband services, and IPv6 may raise operational and policy challenges to the RIRs, yet at the same time they bring opportunities for increased global cooperation, in a context where distinct regional concerns are represented more effectively than ever before.

It is hoped that the emergence of new RIRs will only serve to expand and enhance the inclusive nature of RIR activities.

References

- [1] Clark, D., and Cohen, D., "A Proposal for Addressing and Routing in the Internet," IEN 46, June 1978.
- [2] Postel, J., "Assigned Numbers," RFC 790, September 1981.
- [3] Information Sciences Institute, "Internet Protocol, DARPA Internet Program, Protocol Specification," RFC 791, September 1981.
- [4] Cerf, V., "IAB Recommended Policy on Distributing Internet Identifier Assignment and IAB Recommended Policy Change to Internet 'Connected' Status," RFC 1174, August 1990.
- [5] Williamson, S., and Nobile, L., "Transition of NIC Services," RFC 1261, September 1991.
- [6] Fuller, V., Li, T., Yu, J., and Varadhan, K., "Supernetting: An Address Assignment and Aggregation Strategy," RFC 1338, June 1992.
- [7] Gerich, E., "Guidelines for Management of IP Address Space," RFC 1366, October 1992.
- [8] Gerich, E., "Guidelines for Management of IP Address Space," RFC 1466, May 1993.
- [9] Rekhter, Y., and Li, T., "A Border Gateway Protocol 4 (BGP-4)," RFC 1654, July 1994.
- [10] Hubbard, K., Kesters, M., Conrad, D., Karrenberg, D., and Postel, J., "Internet Registry IP Guidelines," RFC 2050, November 1996.
- [11] Blokzijl, R., Devillers, Y., Karrenberg, D., and Volk, R., "RIPE Network Coordination Center," RIPE-19, September 1990.
- [12] Terpstra, M., "RIPE NCC Internet Numbers Registration Procedures," RIPE-65, July 1992.

DANIEL KARRENBORG has helped to build the European Internet since the early 1980s. As one of the founding members of the German UNIX Users Group, he has been involved in the setting up of EUnet, a pan-European cooperative network providing electronic mail and news to businesses and academic institutions all over Europe. While at CWI in Amsterdam, Karrenberg helped to expand this network and convert it to a fully IP-based service. During this time he created a whois database of operational contacts, which was the nucleus of the current RIPE database. Karrenberg is one of the founders of RIPE, the IP coordination body for Europe and surrounding areas. In 1992 he was asked to set up the RIPE NCC, the first regional Internet registry providing IP numbers to thousands of Internet service providers in more than 90 countries. Karrenberg led the RIPE NCC until 1999, when it had an international staff of 59 with more than 20 nationalities; he currently helps to develop new RIPE NCC services. Recently his contributions have been recognized by the Internet Society with its *Jon Postel Service Award*. Karrenberg's current interests include measurements of Internet performance and routing as well as security within the Internet infrastructure. In general he likes building new and interesting things. Mr. Karrenberg holds an MSc in computer science from Dortmund University. E-mail: **Daniel.Karrenberg@ripe.net**

GERARD ROSS holds a BA and LLB from University of Queensland and a Grad.Dip. (Communication) from Queensland Institute of Technology. He was employed as the technical writer at APNIC in 1998 and has been involved in the development and drafting of several major policy documents both in the APNIC region and as part of coordinated global RIR activities. He was the ASO webmaster in its inaugural year. He is currently the APNIC Documentation Manager. E-mail: **gerard@apnic.net**

PAUL WILSON has been Director-General of APNIC since August 1998. Previously, he was a founding staff member and subsequently Chief Executive Officer at Pegasus Networks, the first private ISP in Australia. Over an eight-year period he worked as a consultant to the United Nations and other international agencies on Internet projects in many countries. Since 1994, he has worked with the International Development Research Centre (IDRC) on its Pan-Asia Networking (PAN) Programme, supporting projects in Mongolia, Vietnam, Cambodia, Maldives, Nepal, Bhutan, PNG, and China. He continues to serve as a member of the PAN Research and Development Grants Committee. E-mail: **pwilson@apnic.net**

LESLIE NOBILE received her B.A. from the American University in Washington, D.C. She has over 15 years of experience in the Internet field, and has been involved with the Internet Registry system since 1991. Prior to that, she held various technical management positions while working under a U.S. Government contract that supported the engineering and implementation of the Defense Data Network, a high-speed data network that evolved from the ARPANET. Her experience with the Registry system began in 1991 working as one of the Operations managers who transitioned the Internet Network Information Center (NIC) from SRI to Network Solutions, Inc. She remained a registration services manager with the DDN/DoD NIC until August 2000, when she became Director of Registration Services at the American Registry for Internet Numbers (ARIN). She has been a contributing author to RFCs, Internet Society (ISOC) articles, and various other industry publications and has been actively involved in the global coordination of Internet addressing policy. Her e-mail address is **leslie@arin.net**

Book Reviews

Web Caching *Web Caching* by Duane Wessels, ISBN 1-56592-536-X, O'Reilly, June 2001.

It's always a pleasure to read a technical book written by someone who has not just studied the topic, but has been so involved that he has spent years living and breathing the subject. Such books do more than just describe the technology, because they are invariably able to add a dimension of deeper insight and interest, and in so doing, bring the topic to life for the reader. Duane Wessel's experiences in the Harvest project, and then as self-confessed "Chief Procrastinator" in the *Squid* Web cache project, certainly place him in the category of an author who has lived the topic. The outcome is a well-researched and very readable book on the topic of Web caching.

Web Caching

Web caching has been an integral part of the architecture of the World Wide Web since its inception, and is now a broad topic encompassing a range of approaches, a range of technologies, and a range of deployment issues for the end consumer, the content publisher, and the service provider intermediaries. The book starts with a clear introduction that outlines the elements of the architecture of the Web, and describes the terminology used within the book. This section also provides a basic introduction to the operation of the *Hypertext Transfer Protocol* (HTTP). This section also describes the various forms of Web caches that are in use today.

The way in which a cache interprets the directives at the header of a delivered Web object is described in some detail. I learned something unexpected here, in that a Web object that includes a directive of the form "Cache-control: no-cache" is defined in RFC 2616 as allowing a cache to store a copy of the object and use it, subject to revalidation, for subsequent requests. It seems that if you really want the object not to be stored in a cache, then "no-store" is what you are after, because "no-cache" allows the object to be cached! As well as describing the definition of the cache control directives, this section provides a clear explanation of how document ageing is defined, and when a cache server determines that a cached object should be checked against the original to ensure that the cached copy remains a faithful reproduction.

Caching has its champions and its detractors, and the book attempts to present both perspectives in a balanced fashion. On the positive side, caching is seen as an effective way to improve the performance of the delivery of Web-based services, and to relieve network and server load. The claim is made here that a large busy cache can achieve a hit ratio of some 70 percent. Don't get too enthusiastic, however, because a more common achieved ratio is somewhere between 30 and 40 percent.

On the negative side is the ever-present issue of accuracy of the cache, the inability for a content provider to track content access, and the issue of integrity of the cache in the face of service attacks that are directed to the cached copy of the content.

The Politics of Caching

This section of the book intrigued me, because it is certainly rare to see a technical book address the various social implications of the technology. The study includes the issues of privacy, request blocking, copyright control, content integrity, cache busting, and the modifications to the trust model in the presence of cache intermediaries. The book exposes the tension between the content provider, the user, and the service provider. The content provider would generally like to exercise some control over tracking who is accessing the content and how each client uses the content and how they navigate through the Web site. The user is interested in efficiency of content delivery, and also has to place a high level of trust in the integrity of the content-delivery system. The service provider is also interested in rapid delivery of content, as well as managing network load. Third parties, such as regulatory or law-enforcement bodies, may be interested in ensuring that the content originator is unambiguously traceable, and that various regulations with respect to content are enforced by content originators and service providers.

Practical Advice

From this overview, the book moves onto more practical topics, and first describes how to configure browsers to take advantage of caches. It also covers how various proxy auto-configurators work. The topic that has generated some attention is that of *interception caching*, where a user's Web-browser commands are intercepted by a provider cache without the direct knowledge of the user of the user's browser. The techniques of implementing such interception caches are described, including a description of the operation of the *Web Cache Coordination Protocol* (WCCP), policy routing, and firewall interception. Interception caching, or transparent caching, is a topic that has generated its fair share of controversy in the past, and the book does take the time to clearly describe the issues associated with this caching approach.

The other topic covered under the general topic of practical advice is advice to server operators and content providers on how to make servers and content work in a predictable fashion with caches, describing which HTTP reply headers affect cacheability. This section provides advice on how to build a cache-friendly Web site, and motivates this with reasons why a content provider would want to ensure that content is readily cacheable. This includes some practical advice on how a content provider can still receive hit counts and site navigation information while still allowing the content of a site to be cached.

Fun with Caches—Cache Hierarchies and Clusters

Although caches can operate in a standalone configuration, it is possible to interconnect caches so that a cache will refer to another cache in the event of a cache miss, rather than directly refer to the origin server. I gather that the author is not overly keen on such an approach, given that the arguments against such configurations consume five times as much space as the arguments in favor! The alternative to a strict hierarchy is a set of cooperating peer caches, together with an intercache protocol to allow a cache to efficiently query its peers for an object. The book describes the *Internet Cache Protocol* (ICP), the *Cache Array Routing Protocol* (CARP), which is pointed out to be an algorithm, not a protocol, despite its name, the *Hypertext Caching Protocol* (HTCP), and *Cache Digests*. The scenarios where each approach would be preferred is a helpful addition to this section. Cache clusters are also described; if I have a criticism of the book, it is that this section is too terse—I was looking for more details of cache-balancing and content-distribution techniques.

Cache Operation

The final section of the book looks at the tasks associated with designing, benchmarking, and operating cache servers. How much disk space is enough for a cache? How much memory? Where should the caches be placed in the network? What aspects of the cache operation should you monitor? And if you are considering purchasing caches, what aspects of the cache should you carefully examine?

Conclusion

This is not a book about how to build a cache, although if you are considering doing that it's a good place to start your research. Nor is it a book about every detail on how to operate a cache. But if you are operating a cache, it will be useful. Although it's not a book about how to operate a Web server, if you are operating a Web server, then caches will attempt to store your content, and this book will help you configure your server to interoperate predictably with caches.

The Web is a large part of today's Internet, and Web caches can make the Web faster, more efficient, and more resilient. If you want to understand how caches work and understand how you can use caches to improve the user's experience rather than making things worse, then this book is essential reading.

—Geoff Huston
gih@telstra.net

IPSec *IPSec: The New Security Standard for the Internet, Intranets, and Virtual Private Networks*, by Naganand Doraswamy and Dan Harkins, ISBN 0-13-011898-2, 1999, Prentice Hall PTR Web Infrastructure series. <http://www.phptr.com>

We all know that Internet security is a major concern. Evolving technologies such as *Virtual Private Networks* (VPNs) are making it easier to deploy secure networks at low costs. VPN technology is based upon encryption techniques that make use of different algorithms. Most of these algorithms are specified in the form of *Requests for Comments* (RFCs). Though RFCs provide the minute details, they are not exactly lively reading. This is where the *IP Security* (IPSec) book comes in handy. The authors have done their best to explain IPSec technology in layman's language, although one encounters a lot of technical jargon in this book.

Organization

The book is divided into three parts. Part I gives a history of cryptography and techniques and cryptographic tools, and overviews of TCP/IP and IPSec. Authentication methods such as *Public Key Infrastructure* (PSI), RSA, and DSA are discussed. Key exchange methods such as Diffie-Hellman and RSA Key Exchange are discussed, along with their advantages and disadvantages. IPSec architecture is explored in the IP Security Overview section, which describes the security services provided by IPSec, how packets are constructed and processed, and the interaction of IPSec processing with policy. IPSec protocols—*Authentication Header* (AH) and *Encapsulation Security Payload* (ESP)—are the basic ingredients of the IPSec stack to provide security. Both AH and ESP can be operated in either the transport mode or tunnel mode. Part II offers a detailed analysis of IPSec, the different modes, IPSec implementation, the ESP, AH, and the *Internet Key Exchange* (IKE). The authors do a good job of describing the IPSec road map, which defines how various components within IPSec interact with each other. Detailed packet formats of different IPSec formats are discussed in Chapter 4. ESP, AH, and IKE are discussed in depth in Chapters 5 through 7. Part III deals with most of the deployment issues concerned with IPSec, as well as policy definition, policy management, implementation architecture, and end-to-end security are discussed in this section. Chapter 11 discusses the future of IPSec and what it means to the world of security. Though IPSec may be thought of as a totally secure method of communication, it has its conflicts when it comes to *Network Address Translation* (NAT), multicasting, and key management in a multicast environment.

Prerequisites

Although the authors have done a good job delivering the IPSec concept, understanding this text requires more than basic computer and communication concepts. One should understand hacking and different types of Internet attacks. OSI layer details and packet-level understanding of every layer within the OSI model is a must.

—Manohar Chandrashekar, WorldCom Inc

mchandra@wcom.com

Letters to the Editor

ICANN Mr. Jacobsen,

I very much enjoy the *Internet Protocol Journal* and put it at the top of my reading stack as soon as it is received. In particular, I enjoy the standards and high technical detail and view it as a safe place from overt commercial advertisement and politics.

That is why I was disappointed by the article from Mr. Lynn. My opinion of ICANN is that it is undemocratic in any tradition, uninterested in experimentation, and uninterested in outside views. I took offense at his continued use of the phrase “public trust” and interpreted the article as propaganda. Further, I found the technical content of the article to be zero.

On the other hand, William Stallings article on MPLS was exactly the kind of article I’ve come to enjoy. I wasn’t familiar with MPLS and the article helped me understand the concepts, vocabulary, and high-level issues. I hope that “MPLS” serves as a model of the articles in future IPJ issues.

I keep back issues of IPJ in a binder and continue to hope you uncover more articles like “The Social Life of Routers.” My copy of Mr. Krebs article has notes in all the margins—I was excited—but it was a twist on something that I thought I knew and he exposed a different design vocabulary by making an unexpected comparison.

I apologize for complaining about something that is a gift from Cisco; I do understand how crass that is. I hope that you will interpret my note in a complementary manner: I’ve come to respect the journal and found that it fits an unfilled niche in my reading.

—Brent D. Stewart, Global Knowledge
<brent@stewart.hickory.nc.us>

Brent,

I appreciate your feedback, as I am sure Mr. Lynn will if you send it to him. The article was, after all, published for public comment.

ICANN has unfortunately tended to polarize people and has become a forum in which a certain amount of politics is played out. I don’t think this is entirely ICANN (the board)’s fault. What was set up as an organization to take over the work of one man—the late Jon Postel, is seen by some as an opportunity for “Internet Governance” and “world-wide electronic democracy.”

Having watched the ICANN process since its beginnings in 1998, I would say that Mr. Lynn’s version of history is pretty much on target. When the IANA was in the hands of Jon Postel, it most certainly was a “public trust” (a limited resource to say the least), and if ICANN does not take that responsibility seriously, it certainly will have failed.

However, I do not think this is the case. Yes, ICANN is now a fairly large and slow moving machinery, and I would have liked to see more new domains deployed sooner, but to some extent the slowness is caused by the structure of Supporting Organizations as much as it is by the board itself. There is a lot to sort out, a lot to comment on, and *many* divergent views are indeed being expressed in all kinds of ICANN forums, including the public meetings. So, I cannot agree that ICANN is “uninterested in outside views.” A perfect democracy it is not, nor was it ever intended to be, and yes, some of the topics on the agenda such as the *Uniform Dispute Resolution Process* (UDRP) are indeed non-technical. But it is not as if ICANN had much choice in that particular matter. (Although some would argue that it could be moved outside the ICANN process.)

Being part of the ICANN process, through e-mail discussion, public meetings or through the Supporting Organizations is not difficult. Nor do I think that ICANN ignores any of the feedback it gets.

Back to the article. No, it was not particularly technical, but if you read IPJ’s Call for Papers you will see that it mentions “Legal, policy and regulatory topics...” Also, in the wake of September 11, I thought it was important to provide some background on the thinking of ICANN, and why they chose to refocus the most recent meeting on security etc. IPJ, by the way, also encourages the occasional “Opinion Piece,” although the article by Mr. Lynn was not intended as such. The issue of alternate roots is indeed a matter of debate, and while the IAB has already expressed its view, I appreciate that there might be other (valid) ones.

In any case, thank you for taking the time to write. I certainly don’t intend to steer IPJ away from topics such as MPLS and I hope that the occasional policy or even opinion piece won’t steer you away from IPJ.

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

MPLS Ole,

William Stallings otherwise-excellent article on MPLS in the *Internet Protocol Journal* Vol. 4, No. 3 had a serious error in it with respect to Virtual Private Networks (VPNs). He said that MPLS is an efficient mechanism for supporting VPNs and that MPLS provides security; neither is true.

As the rest of the article shows, MPLS provides a transport tunnel for IP packets, meaning that it helps create virtual networks. However, there is no privacy on those virtual networks, so it is inappropriate and probably dangerous to call MPLS tunnels virtual private networks.

To most Internet users, security means preventing snooping of sensitive traffic, preventing malicious changes to content, or both. MPLS does not provide either service. Instead of relying on insecure MPLS, users who want secure tunnels use systems that employ the IPsec protocol.

Many dozens of vendors supply IPsec systems appropriate for everything from tiny home offices to gigantic telco central switches, all with the same high security. Although the article showed that MPLS has many valuable features, IPJ readers should not fall into the trap of thinking that VPN support or security are MPLS features.

—Paul Hoffman, Director, VPN Consortium
<paul.hoffman@vpnc.org>

Ed: We presented this letter to a panel of experts, and here are some samples of the responses we received:

The term “VPN” has been used in many different contexts. I saw a group once call a VLAN a “VPN” as well. I honestly couldn’t say that they were incorrect. It may be appropriate to say that there are IPsec VPNs and that there are MPLS VPNs, but I have a problem calling one “right” and another “wrong” simply because of some perceived, implied definition of the security level that should be provided by a “VPN.” Most people support the notion that an MPLS VPN provides about as much “security” as a Frame Relay link. This amount of “security” in a VPN is acceptable to many people.

—Chris Lowick, Cisco Systems <clonvick@cisco.com>

We have different views on security, I’m sure. One view is that a secure private network: a) ensures that a third party cannot impose a condition on the network such that a customer’s traffic is directed to another customer b) ensures that a third party cannot inject traffic into a customer’s private network, c) a third party cannot alter customer traffic and d) a third party cannot discern that communications is taking place between two parts of a private network.

MPLS uses the same mechanisms as X.25, ATM and Frame, and has similar properties—the objectives above can be met with adequate confidence as long as the network is carefully configured and managed.

Edge to edge IPsec has a different set of security principles—the basic mode of operation is that such networks may be subject to attacks that redirect customer’s traffic to third party sites, and allow third parties to inject traffic into the VPN, and allow a third party to discern that communications is taking place within a private context. The essential attribute of edge to edge IPsec is that the encryption is intended to ensure that leakage can be identified: foreign injected traffic or altered traffic can be identified and rejected and leaking traffic cannot be decoded.

Both approaches have vulnerabilities and weaknesses. The first approach places trust in the integrity of the host platform. The second approach is prone to various forms of DOS attacks and traffic profiling.

But I would not concur with a view that labels the MPLS approach as inefficient or insecure, nor would I label X.25 networks, ATM or Frame as *intrinsically* inefficient and insecure. There are insecure operating practices and there are cautious operating practices.

IPSec networks have similar issues—relating particularly to the vulnerabilities of third party disruption and profiling eavesdropping.

So it's not that I believe that all MPLS networks are well designed and well operated—on the contrary! But as an architectural approach I am not able to agree with a comment that appears to condemn MPLS as *intrinsically* a poor choice for a VPN host technology.

So if the comment is that the article provides the impression that MPLS is such a robust technology that it creates secure private network applications such as VPNs, and appears to make this assertion so strongly that it gives the impression that this outcome occurs irrespective of MPLS network design and operating practices, and that this impression is ill-founded, then I would agree entirely with Mr. Hoffman. Secure networks, or at least robust networks, are a result of careful choice of technologies coupled with careful design and careful operation.

—Geoff Huston, Telstra <gih@telstra.net>

Ed.: We forwarded these comments to Mr. Hoffman, and he responded:

Geoff believes that it a network that does not prevent an active attacker from seeing or modifying traffic, and does not prevent a passive attacker from seeing packets, is secure and private; I do not. The fact that MPLS restricts the flow of traffic to a particular defined network is sufficient for him; it is not for me, given the fact that an attacker breaking into any node on that defined network can compromise the privacy and integrity of the traffic.

It is typical for ISPs to not want to do the work of actually securing the traffic they say they have put in a VPN by using IPsec. That work is not cheap, and takes more management than vanilla MPLS, but it is the only way to really secure the data. I am absolutely not saying that the IPsec community is without blame here: we have a tendency to ignore the valuable features of MPLS and have done almost nothing to make it easier to intelligently tunnel IPsec in MPLS (we also pretty much stonewalled the IPsec under L2TP work that is now finally standardized). But our lack of openness doesn't make MPLS a VPN technology.

—Paul Hoffman, Director, VPN Consortium
<paul.hoffman@vpnc.org>

Ed.: We would love to hear from you. Please send your letters to:
ipj@cisco.com

ACM Assembles Security and Privacy Panel

Prompted by increased public concerns about personal privacy and the security of networked information systems, the *Association for Computing Machinery* (ACM) has announced the formation of a new *Advisory Committee on Security and Privacy* (ACSP). Led by Peter Neumann and Eugene H. Spafford, the ACSP brings together a dozen leaders and innovators in the field of privacy and information assurance to serve as a powerful resource for the ACM community and the public at large.

Comprising experts from research, industry, academia, and government, the diverse group represents a wide range of viewpoints. Commenting on the formation of the ACSP, Co-Chair Peter Neumann noted, “The ACSP will provide timely and accurate assessments of situations relating to information security that are otherwise clouded by confusion, uncertainty, and often, misinformation.”

Added ACSP Co-Chair Gene Spafford, “Until recently, computing professionals have been primarily concerned with making computers work consistently, cheaply, and effectively. Now it is critical that we also bring expertise to bear on how computers can be made to operate safely, keep information resources secure from attack, and protect privacy.”

The ACSP consists of 12 distinguished members with expertise in information security and assurance, privacy, cybercrime, and allied fields. The group will coordinate with other ACM Committees, including the *U.S. ACM Committee on Public Policy* (USACM) and ACM Law Committee, to provide objective advice to the computing community, the public at large, and to policy-makers. ACSP is expected to provide statements and testimony on information security and privacy issues, as well as undertaking studies of related topics. For more information about the ACSP, see the web site at:

<http://www.acm.org/usacm/ACSP/homepage.htm>

Members of the ACSP (affiliations provided for identification purposes only) are:

Steve Bellovin (AT&T Labs Research)
Matthew Blaze (AT&T Labs Research)
David Clark (MIT)
Dorothy Denning (Georgetown University)
Ed Felten (Princeton University)
David Farber (University of Pennsylvania)
Susan Landau (Sun Microsystems)
Robert Morris (Dartmouth College)
Peter Neumann (SRI International)
Fred Schneider (Cornell University)
Eugene H. Spafford (Purdue University CERIAS)
Willis Ware (RAND Corporation)

For more information, see ACM’s Web site at: **<http://www.acm.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, trouble-shooting, 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 will contain 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

Ole J. Jacobsen, Editor and Publisher

Editorial Advisory Board

Dr. Vint Cerf, Sr. VP, Internet Architecture and Technology
WorldCom, USA

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

David Farber
The Alfred Fitler Moore Professor of Telecommunication Systems
University of Pennsylvania, USA

Peter Löthberg, Network Architect
Stupi AB, Sweden

Dr. Jun Murai, Professor, WIDE Project
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*

*Cisco, Cisco Systems, and the Cisco
Systems logo are registered
trademarks of Cisco Systems, Inc. in
the USA and certain other countries.
All other trademarks mentioned in this
document are the property of their
respective owners.*

*Copyright © 2001 Cisco Systems Inc.
All rights reserved. Printed in the USA.*



The Internet Protocol Journal, Cisco Systems
170 West Tasman Drive, M/S SJ-10/5
San Jose, CA 95134-1706
USA

ADDRESS SERVICE REQUESTED

PRSR STD
U.S. Postage
PAID
Cisco Systems, Inc.