

Network Working Group
Request for Comments: 4029
Category: Informational

M. Lind
TeliaSonera
V. Ksinant
Thales Communications
S. Park
SAMSUNG Electronics
A. Baudot
France Telecom
P. Savola
CSC/Funet
March 2005

Scenarios and Analysis for Introducing IPv6 into ISP Networks

Status of This Memo

This memo provides information for the Internet community. It does not specify an Internet standard of any kind. Distribution of this memo is unlimited.

Copyright Notice

Copyright (C) The Internet Society (2005).

Abstract

This document describes different scenarios for the introduction of IPv6 into an ISP's existing IPv4 network without disrupting the IPv4 service. The scenarios for introducing IPv6 are analyzed, and the relevance of already defined transition mechanisms are evaluated. Known challenges are also identified.

Table of Contents

1.	Introduction.	2
1.1.	Goal and Scope of the Document.	2
2.	Brief Description of a Generic ISP Network.	3
3.	Transition Scenarios.	4
3.1.	Identification of Stages and Scenarios.	4
3.2.	Stages.	5
3.2.1.	Stage 1 Scenarios: Launch	5
3.2.2.	Stage 2a Scenarios: Backbone.	6
3.2.3.	Stage 2b Scenarios: Customer Connection	6
3.2.4.	Stage 3 Scenarios: Complete	7
3.2.5.	Stages 2a and 3: Combination Scenarios.	7
3.3.	Transition Scenarios.	7
3.4.	Actions Needed When Deploying IPv6 in an ISP's Network.	8

4.	Backbone Transition Actions	9
4.1.	Steps in the Transition of Backbone Networks.	9
4.1.1.	MPLS Backbone	9
4.2.	Configuration of Backbone Equipment	10
4.3.	Routing	10
4.3.1.	IGP	11
4.3.2.	EGP	12
4.3.3.	Transport of Routing Protocols.	12
4.4.	Multicast	13
5.	Customer Connection Transition Actions.	13
5.1.	Steps in the Transition of Customer Connection Networks	13
5.1.1.	Small End Sites	14
5.1.2.	Large End Sites	15
5.2.	User Authentication/Access Control Requirements	15
5.3.	Configuration of Customer Equipment	16
5.4.	Requirements for Traceability	16
5.5.	Ingress Filtering in the Customer Connection Network.	17
5.6.	Multihoming	17
5.7.	Quality of Service.	17
6.	Network and Service Operation Actions	18
7.	Future Stages	18
8.	Requirements for Follow-On Work	19
9.	Example Networks.	19
9.1.	Example 1	21
9.2.	Example 2	22
9.3.	Example 3	23
10.	Security Considerations	23
11.	Acknowledgments	24
12.	Informative References.	24
	Appendix A.	26
	Authors' Addresses.	27
	Full Copyright Statement.	28

1. Introduction

1.1. Goal and Scope of the Document

When an ISP deploys IPv6, its goal is to provide IPv6 connectivity and global address space to its customers. The new IPv6 service must be added to an existing IPv4 service, and the introduction of IPv6 must not interrupt this IPv4 service.

An ISP offering IPv4 service will find different ways to add IPv6 to this service. This document discusses a small set of scenarios for the introduction of IPv6 into an ISP's IPv4 network. It evaluates the relevance of the existing transition mechanisms in the context of these deployment scenarios and points out the lack of essential functionality in these methods.

The document is focused on services that include both IPv6 and IPv4 and does not cover issues surrounding IPv6-only service. It is also outside the scope of this document to describe different types of access or network technologies.

2. Brief Description of a Generic ISP Network

A generic network topology for an ISP can be divided into two main parts: the backbone network and customer connection networks. In addition, it includes building blocks such as network and service operations. The additional building blocks used in this document are defined as follows:

"CPE" : Customer Premises Equipment

"PE" : Provider Edge Equipment

"Network and service operation"

: This is the part of the ISP's network that hosts the services required for the correct operation of the ISP's network. These services usually include management, supervision, accounting, billing, and customer management applications.

"Customer connection"

: This is the part of the network used by a customer when connecting to an ISP's network. It includes the CPE, the last hop link, and the parts of the PE interfacing to the last hop link.

"Backbone" : This is the rest of the ISP's network infrastructure. It includes the parts of the PE interfacing to the core, the core routers of the ISP, and the border routers used to exchange routing information with other ISPs (or other administrative entities).

"Dual-stack network"

: A network that natively supports both IPv4 and IPv6.

In some cases (e.g., incumbent national or regional operators), a given customer connection network may have to be shared between or among different ISPs. According to the type of customer connection network used (e.g., one involving only layer 2 devices or one involving non-IP technology), this constraint may result in architectural considerations relevant to this document.

The basic components in the ISP's network are depicted in Figure 1.

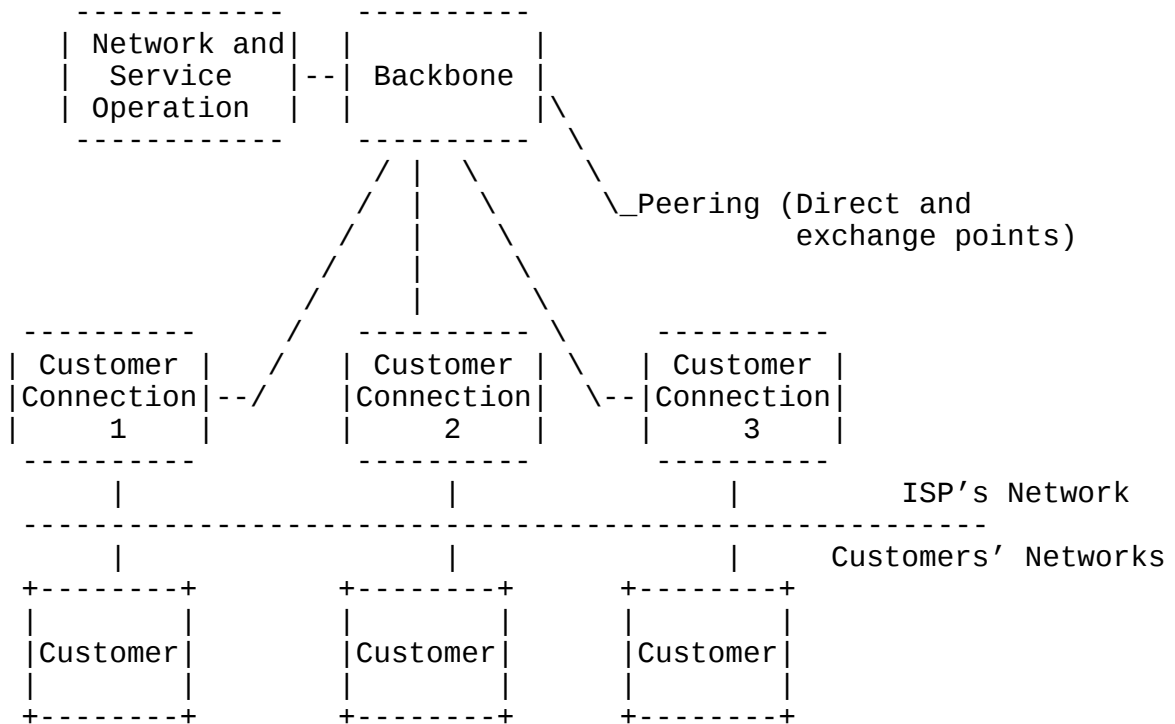


Figure 1: ISP Network Topology

3. Transition Scenarios

3.1. Identification of Stages and Scenarios

This section describes different stages an ISP might consider when introducing IPv6 connectivity into its existing IPv4 network and the different scenarios of what might occur in the respective stages.

The stages here are snapshots of the ISP's network with respect to IPv6 maturity. Because the ISP's network is continually evolving, a stage is a measure of how far along the ISP has come in terms of implementing the functionality necessary to offer IPv6 to its customers.

It is possible for a transition to occur freely between different stages. Although a network segment can only be in one stage at a time, the ISP's network as a whole can be in different stages. Different transition paths can be followed from the first to the final stage. The transition between two stages does not have to be instantaneous; it can occur gradually.

Each stage has different IPv6 properties. Therefore, based on its requirements, an ISP can decide which set of stages it will follow and in what order to transform its network.

This document is not aimed at covering small ISPs, hosting providers, or data centers; only the scenarios applicable to ISPs eligible for at least a /32 IPv6 prefix allocation from an RIR are covered.

3.2. Stages

The stages are derived from the generic description of an ISP's network in Section 2. Combinations of different building blocks that constitute an ISP's environment lead to a number of scenarios from which the ISP can choose. The scenarios most relevant to this document are those that maximize an ISP's ability to offer IPv6 to its customers in the most efficient and feasible way. The assumption in all stages is that the ISP's goal is to offer both IPv4 and IPv6 to the customer.

The four most probable stages are as follows:

- o Stage 1 Launch
- o Stage 2a Backbone
- o Stage 2b Customer connection
- o Stage 3 Complete

Generally, an ISP is able to upgrade a current IPv4 network to an IPv4/IPv6 dual-stack network via Stage 2b, but the IPv6 service can also be implemented at a small cost by adding simple tunnel mechanisms to the existing configuration. When a new network is designed, Stage 3 might be the first or last step because there are no legacy concerns. Nevertheless, the absence of IPv6 capability in the network equipment can still be a limiting factor.

Note that in every stage except Stage 1, the ISP can offer both IPv4 and IPv6 services to its customers.

3.2.1. Stage 1 Scenarios: Launch

The first stage is an IPv4-only ISP with an IPv4 customer. This is the most common case today and is the natural starting point for the introduction of IPv6. From this stage, the ISP can move (undergo a transition) from Stage 1 to any other stage with the goal of offering IPv6 to its customer.

The immediate first step consists of obtaining a prefix allocation (typically a /32) from the appropriate RIR (e.g., AfriNIC, APNIC, ARIN, LACNIC, RIPE) according to allocation procedures.

The ISP will also need to establish IPv6 connectivity to its upstream providers and peers; it is of utmost importance to require IPv6 transit when negotiating IP transit deals with the upstream ISPs. If the upstream is not providing IPv6 connectivity at the moment, it may be possible to obtain temporary connectivity from a nearby ISP, possibly using a short configured tunnel. However, the longer-term goal must be to require and to obtain IPv6 connectivity from the transit ISPs, because otherwise the quality of IPv6 connectivity will likely be poor.

Connectivity to peers can typically be established either directly or at Internet Exchange Points (IX). Most IXs use techniques where IPv6 is easy to use, and many IXs already provide infrastructure for IPv6 peerings. Such peerings can be done natively by using IPv6. Peerings over IPv6-in-IPv4 tunnels is also possible but not recommended, at least in the long term. Direct connectivity to peers may be feasible when there is direct connectivity to the peer for IPv4.

3.2.2. Stage 2a Scenarios: Backbone

Stage 2a deals with an ISP with IPv4-only customer connection networks and a backbone that supports both IPv4 and IPv6. In particular, the ISP has the possibility of making the backbone IPv6-capable through software upgrades, hardware upgrades, or a combination of both.

Since the customer connections have not yet been upgraded, a tunneling mechanism has to be used to provide IPv6 connectivity through the IPv4 customer connection networks. The customer can terminate the tunnel at the CPE (if it has IPv6 support) or at some set of devices internal to its network. That is, either the CPE or a device inside the network could provide global IPv6 connectivity to the rest of the devices in the customer's network.

3.2.3. Stage 2b Scenarios: Customer Connection

Stage 2b consists of an ISP with an IPv4 backbone network and a customer connection network that supports both IPv4 and IPv6. Because the service to the customer is native IPv6, the customer is not required to support both IPv4 and IPv6. This is the biggest difference from the previous stage. The need to exchange IPv6 traffic still exists but might be more complicated than in the previous case because the backbone is not IPv6-enabled. After completing Stage 2b, the original IPv4 backbone is unchanged. This means that the IPv6 traffic is transported either by tunneling over the existing IPv4 backbone, or in an IPv6 overlay network more or less separated from the IPv4 backbone.

Normally, the ISP will continue to provide IPv4 connectivity by using private (NATted by the ISP) or public IPv4 address. In many cases, the customer also has a NAT of his/her own; if so, this likely continues to be used for IPv4 connectivity.

3.2.4. Stage 3 Scenarios: Complete

Stage 3 could be considered the final step in introducing IPv6, at least within the scope of this document. This stage consists of ubiquitous IPv6 service with native support for IPv6 and IPv4 in both backbone and customer connection networks. From the customer's perspective, it is identical to the previous stage because the customer connection network has not changed. The requirement for exchanging IPv6 traffic is identical to that of Stage 2.

3.2.5. Stages 2a and 3: Combination Scenarios

Some ISPs may use different access technologies of varying IPv6 maturity. This may result in a combination of the Stages 2a and 3: some customer connections do not support IPv6, but others do; in both cases the backbone is dual-stack.

This scenario is equivalent to Stage 2a, but it requires support for native IPv6 customer connections on some access technologies.

3.3. Transition Scenarios

Given the different stages, it is clear that an ISP has to be able to make a transition from one stage to another. The initial stage in this document is an IPv4-only service and network. The end stage is a dual IPv4/IPv6 service and network.

The transition starts with an IPv4 ISP and then moves in one of three directions. This choice corresponds to the different transition scenarios. Stage 2a consists of upgrading the backbone first. Stage 2b consists of upgrading the customer connection network. Finally, Stage 3 consists of introducing IPv6 in both the backbone and customer connections as needed.

Because most ISP backbone IPv4 networks continually evolve (firmware replacements in routers, new routers, etc.), they can be made ready for IPv6 without additional investment (except staff training). This transition path may be slower but still useful, as it allows for the introduction of IPv6 without any actual customer demand. This approach may be superior to doing everything at the last minute, which may entail a higher investment. However, it is important to consider (and to request from vendors) IPv6 features in all new equipment from the outset. Otherwise, the time and effort required

to remove non-IPv6-capable hardware from the network may be significant.

3.4. Actions Needed When Deploying IPv6 in an ISP's Network

Examination of the transitions described above reveals that it is possible to split the work required for each transition into a small set of actions. Each action is largely independent of the others, and some actions may be common to multiple transitions.

Analysis of the possible transitions leads to a small list of actions:

- * Actions required for backbone transition:
 - Connect dual-stack customer connection networks to other IPv6 networks through an IPv4 backbone.
 - Transform an IPv4 backbone into a dual-stack one. This action can be performed directly or through intermediate steps.
- * Actions required for customer connection transition:
 - Connect IPv6 customers to an IPv6 backbone through an IPv4 network.
 - Transform an IPv4 customer connection network into a dual-stack one.
- * Actions required for network and service operation transition:
 - Set up IPv6 connectivity to upstream providers and peers.
 - Configure IPv6 functions into network components.
 - Upgrade regular network management and monitoring applications to take IPv6 into account.
 - Extend customer management (e.g., RADIUS) mechanisms to be able to supply IPv6 prefixes and other information to customers.
 - Enhance accounting, billing, and so on to work with IPv6 as needed. (Note: If dual-stack service is offered, this may not be necessary.)
 - Implement security for network and service operation.

Sections 4, 5, and 6 contain detailed descriptions of each action.

4. Backbone Transition Actions

4.1. Steps in the Transition of Backbone Networks

In terms of physical equipment, backbone networks mainly consist of high-speed core and edge routers. Border routers provide peering with other providers. Filtering, routing policy, and policing functions are generally managed on border routers.

In the beginning, an ISP has an IPv4-only backbone. In the end, the backbone is completely dual-stack. In between, intermediate steps may be identified:

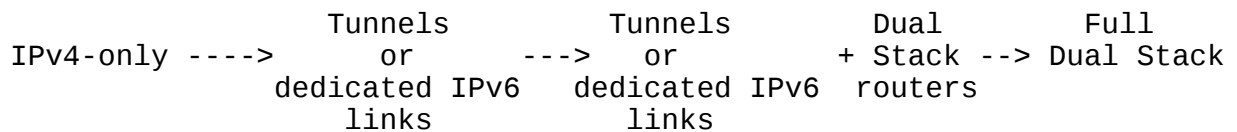


Figure 2: Transition Path

The first step involves tunnels or dedicated links but leaves existing routers unchanged. Only a small set of routers then have IPv6 capabilities. The use of configured tunnels is adequate during this step.

In the second step, some dual-stack routers are added, progressively, to this network.

The final step is reached when all or almost all routers are dual-stack.

For many reasons (technical, financial, etc.), the ISP may progress step by step or jump directly to the final one. One important criterion in planning this evolution is the number of IPv6 customers the ISP expects during its initial deployments. If few customers connect to the original IPv6 infrastructure, then the ISP is likely to remain in the initial steps for a long time.

In short, each intermediate step is possible, but none is mandatory.

4.1.1. MPLS Backbone

If MPLS is already deployed in the backbone, it may be desirable to provide IPv6-over-MPLS connectivity. However, setting up an IPv6 Label Switched Path (LSP) requires signaling through the MPLS

network; both LDP and RSVP-TE can set up IPv6 LSPs, but this might require upgrade/change in the MPLS core network.

An alternative approach is to use BGP for signaling or to perform; for example, IPv6-over-IPv4/MPLS, as described in [BGPTUNNEL]. Some possibilities are preferable to others, depending on the specific environment under consideration. The approaches seem to be as follows:

- 1) Require that MPLS networks deploy native IPv6 routing and forwarding support.
- 2) Require that MPLS networks support native routing and setting up of IPv6 LSPs, used for IPv6 connectivity.
- 3) Use only configured tunneling over IPv4 LSPs.
- 4) Use [BGPTUNNEL] to perform IPv6-over-IPv4/MPLS encapsulation for IPv6 connectivity.

Approaches 1) and 2) are clearly the best target approaches. However, approach 1) may not be possible if the ISP is not willing to add IPv6 support in the network, or if the installed equipment is not capable of high performance native IPv6 forwarding. Approach 2) may not be possible if the ISP is unwilling or unable to add IPv6 LSP set-up support in the MPLS control plane.

Approach 4) can be used as an interim mechanism when other options are unfeasible or undesirable for the reasons discussed above.

Approach 3) is roughly equivalent to approach 4) except that it does not require additional mechanisms but may lack scalability in the larger networks, especially if IPv6 is widely deployed.

4.2. Configuration of Backbone Equipment

In the backbone, the number of devices is small, and IPv6 configuration mainly deals with routing protocol parameters, interface addresses, loop-back addresses, access control lists, and so on.

These IPv6 parameters need to be configured manually.

4.3. Routing

ISPs need routing protocols to advertise reachability and to find the shortest working paths, both internally and externally.

Either OSPFv2 or IS-IS is typically used as the IPv4 IGP. RIPv2 is not usually used in service provider networks, as OSPF and IS-IS are superior IGPs. BGP is the only IPv4 EGP. Static routes also are used in both cases.

Note that it is possible to configure a given network so that it has an IPv6 topology different from its IPv4 topology. For example, some links or interfaces may be dedicated to IPv4-only or IPv6-only traffic, or some routers may be dual-stack whereas others may be IPv4- or IPv6-only. In this case, routing protocols must be able to understand and cope with multiple topologies.

4.3.1. IGP

Once the IPv6 topology has been determined, the choice of IPv6 IGP must be made: either OSPFv3 or IS-IS for IPv6. RIPng is not appropriate in most contexts, due to RIPv2 not being appropriate for IPv4 either, and is therefore not discussed here. The IGP typically includes the routers' point-to-point and loop-back addresses.

The most important decision is whether one wishes to have separate routing protocol processes for IPv4 and IPv6. Separating them requires more memory and CPU for route calculations, e.g., when the links flap. But separation provides a measure of assurance that should problems arise with IPv6 routing, they will not affect the IPv4 routing protocol. In the initial phases, if it is uncertain whether joint IPv4-IPv6 networking is working as intended, running separate processes may be desirable and easier to manage.

The possible combinations are as follows:

- With separate processes:
 - o OSPFv2 for IPv4, IS-IS for IPv6 (only)
 - o OSPFv2 for IPv4, OSPFv3 for IPv6, or
 - o IS-IS for IPv4, OSPFv3 for IPv6
- With the same process:
 - o IS-IS for both IPv4 and IPv6

Note that if IS-IS is used for both IPv4 and IPv6, the IPv4/IPv6 topologies must be "convex", unless the multiple-topology IS-IS extensions [MTISIS] have been implemented (using IS-IS for only IPv4 or only IPv6 requires no convexity). In simpler networks or with careful planning of IS-IS link costs, it is possible to keep even incongruent IPv4/IPv6 topologies "convex". The convexity problem is explained in more detail with an example in Appendix A.

When deploying full dual-stack in the short-term, using single-topology IS-IS is recommended. This may be particularly applicable for some larger ISPs. In other scenarios, choosing between one or two separate processes often depends on the perceived risk to the IPv4 routing infrastructure, i.e., whether one wishes to keep them separate for the time being. If this is not a factor, using a single process is usually preferable for operational reasons: not having to manage two protocols and topologies.

The IGP is typically only used to carry loopback and point-to-point addresses and doesn't include customer prefixes or external routes. Internal BGP (iBGP), as described in the next section, is most often deployed in all routers (PE and core) to distribute routing information about customer prefixes and external routes.

Some of the simplest devices (e.g., CPE routers) may not implement routing protocols other than RIPng. In some cases, therefore, it may be necessary to run RIPng in addition to one of the above IGPs, at least in a limited fashion, and then, by some mechanism, to redistribute routing information between the routing protocols.

4.3.2. EGP

BGP is used for both internal and external BGP sessions.

BGP with multiprotocol extensions [RFC2858] can be used for IPv6 [RFC2545]. These extensions enable the exchange of IPv6 routing information and the establishment of BGP sessions using TCP over IPv6.

It is possible to use a single BGP session to advertise both IPv4 and IPv6 prefixes between two peers. However, the most common practice today is to use separate BGP sessions.

4.3.3. Transport of Routing Protocols

IPv4 routing information should be carried by IPv4 transport and, similarly, IPv6 routing information by IPv6 for several reasons:

- * IPv6 connectivity may work when IPv4 connectivity is down (or vice-versa).
- * The best route for IPv4 is not always the best one for IPv6.
- * The IPv4 and IPv6 logical topologies may be different because the administrator may want to assign different metrics to a physical link for load balancing or because tunnels may be in use.

4.4. Multicast

Currently, IPv6 multicast is not a major concern for most ISPs. However, some of them are considering deploying it. Multicast is achieved by using the PIM-SM and PIM-SSM protocols. These also work with IPv6.

Information about multicast sources is exchanged by using MSDP in IPv4, but MSDP is intentionally not defined for IPv6. Instead, one should use only PIM-SSM or an alternative mechanism for conveying the information [EMBEDRP].

5. Customer Connection Transition Actions

5.1. Steps in the Transition of Customer Connection Networks

Customer connection networks are generally composed of a small set of PEs connected to a large set of CPEs and may be based on different technologies depending on the customer type or size, as well as the required bandwidth or even quality of service. Small unmanaged connection networks used for public customers usually rely on different technologies (e.g., dial-up or DSL) than the ones used for large customers, which typically run managed networks. Transitioning these infrastructures to IPv6 can be accomplished in several steps, but some ISPs, depending on their perception of the risks, may avoid some of the steps.

Connecting IPv6 customers to an IPv6 backbone through an IPv4 network can be considered a first careful step taken by an ISP to provide IPv6 services to its IPv4 customers. Some ISPs may also choose to provide IPv6 service independently from the regular IPv4 service.

In any case, IPv6 service can be provided by using tunneling techniques. The tunnel may terminate at the CPE corresponding to the IPv4 service or in some other part of the customer's infrastructure (for instance, on IPv6-specific CPE or even on a host).

Several tunneling techniques have already been defined: configured tunnels with tunnel broker, 6to4 [RFC3056], Teredo [TEREDO], and so on. Some of these are based on a specific addressing plan independent of the ISP's allocated prefix(es), while others use a part of the ISP's prefix. In most cases, using the ISP's address space is preferable.

A key factor is the presence or absence of NATs between the two tunnel end-points. In most cases, 6to4 and ISATAP are incompatible with NATs, and UDP encapsulation for configured tunnels has not been specified.

Dynamic and non-permanent IPv4 address allocation is another factor a tunneling technique may have to deal with. In this case, the tunneling techniques may be more difficult to deploy at the ISP's end, especially if a protocol including authentication (like PPP for IPv6) is not used. This may need to be considered in more detail.

However, NAT traversal can be avoided if the NAT supports forwarding protocol-41 [PROT041] and is configured to do so.

Firewalls in the path can also break tunnels of these types. The administrator of the firewall needs to create a hole for the tunnel. This is usually manageable, as long as the firewall is controlled by either the customer or the ISP, which is almost always the case.

When the CPE is performing NAT or firewall functions, terminating the tunnels directly at the CPE typically simplifies the scenario considerably, avoiding the NAT and firewall traversal. If such an approach is adopted, the CPE has to support the tunneling mechanism used, or be upgraded to do so.

5.1.1. Small End Sites

Tunneling considerations for small end sites are discussed in [UNMANEVA]. These identify solutions relevant to the first category of unmanaged networks. The tunneling requirements applicable in these scenarios are described in [TUNREQS].

The connectivity mechanisms can be categorized as "managed" or "opportunistic". The former consist of native service or a configured tunnel (with or without a tunnel broker); the latter include 6to4 and, e.g., Teredo -- they provide "short-cuts" between nodes using the same mechanisms and are available without contracts with the ISP.

The ISP may offer opportunistic services, mainly a 6to4 relay, especially as a test when no actual service is offered yet. At the later phases, ISPs might also deploy 6to4 relays and Teredo servers (or similar) to optimize their customers' connectivity to 6to4 and Teredo nodes.

Opportunistic services are typically based on techniques that don't use IPv6 addresses from the ISP's allocated prefix(es), and the services have very limited functions to control the origin and the number of customers connected to a given relay.

Most interesting are the managed services. When dual-stack is not an option, a form of tunneling must be used. When configured tunneling is not an option (e.g., due to dynamic IPv4 addressing), some form of

automation has to be used. Basically, the options are either to deploy an L2TP architecture (whereby the customers would run L2TP clients and PPP over it to initiate IPv6 sessions) or to deploy a tunnel configuration service. The prime candidates for tunnel configuration are STEP [STEP] and TSP [TSP], which both also work in the presence of NATs. Neither is analyzed further in this document.

5.1.2. Large End Sites

Large end sites usually have a managed network.

Dual-stack access service is often a possibility, as the customer network is managed (although CPE upgrades may be necessary).

Configured tunnels, as-is, are a good solution when a NAT is not in the way and the IPv4 end-point addresses are static. In this scenario, NAT traversal is not typically required. If fine-grained access control is needed, an authentication protocol needs to be implemented.

Tunnel brokering solutions have been proposed to help facilitate the set-up of a bi-directional tunnel. Such mechanisms are typically unnecessary for large end-sites, as simple configured tunneling or native access can be used instead. However, if such mechanisms would already be deployed, large sites starting to deploy IPv6 might benefit from them in any case.

Teredo is not applicable in this scenario, as it can only provide IPv6 connectivity to a single host, not the whole site. 6to4 is not recommended due to its reliance on the relays and provider-independent address space, which makes it impossible to guarantee the required service quality and manageability large sites typically want.

5.2. User Authentication/Access Control Requirements

User authentication can be used to control who can use the IPv6 connectivity service in the first place or who can access specific IPv6 services (e.g., NNTP servers meant for customers only). The former is described at more length below. The latter can be achieved by ensuring that for all the service-specific IPv4 access lists, there are also equivalent IPv6 access lists.

IPv6-specific user authentication is not always required. An example would be a customer of the IPv4 service automatically having access to the IPv6 service. In this case, the IPv4 access control also provides access to the IPv6 services.

When a provider does not wish to give its IPv4 customers automatic access to IPv6 services, specific IPv6 access control must be performed parallel with the IPv4 access control. This does not imply that different user authentication must be performed for IPv6, but merely that the authentication process may lead to different results for IPv4 and IPv6 access.

Access control traffic may use IPv4 or IPv6 transport. For instance, RADIUS [RFC2865] traffic related to IPv6 service can be transported over IPv4.

5.3. Configuration of Customer Equipment

The customer connection networks are composed of PE and CPE(s). Usually, each PE connects multiple CPE components to the backbone network infrastructure. This number may reach tens of thousands of customers, or more. The configuration of CPE is difficult for the ISP, and it is even more difficult when it must be done remotely. In this context, the use of auto-configuration mechanisms is beneficial, even if manual configuration is still an option.

The parameters that usually need to be provided to customers automatically are as follows:

- The network prefix delegated by the ISP
- The address of the Domain Name System server (DNS)
- Possibly other parameters (e.g., the address of an NTP server)

When user identification is required on the ISP's network, DHCPv6 may be used to provide configurations; otherwise, either DHCPv6 or a stateless mechanism may be used. This is discussed in more detail in [DUAL-ACCESS].

Note that when the customer connection network is shared between the users or the ISPs and is not just a point-to-point link, authenticating the configuration of the parameters (especially prefix delegation) requires further study.

As long as IPv4 service is available alongside IPv6, it is not required to auto configure IPv6 parameters in the CPE, except the prefix, because the IPv4 settings may be used.

5.4. Requirements for Traceability

Most ISPs have some kind of mechanism to trace the origin of traffic in their networks. This also has to be available for IPv6 traffic, meaning that a specific IPv6 address or prefix has to be tied to a

certain customer, or that records must be maintained of which customer had which address or prefix. This also applies to the customers with tunneled connectivity.

This can be done, for example, by mapping a DHCP response to a physical connection and storing the result in a database. It can also be done by assigning a static address or prefix to the customer. A tunnel server could also provide this mapping.

5.5. Ingress Filtering in the Customer Connection Network

Ingress filtering must be deployed toward the customers, everywhere, to ensure traceability, to prevent DoS attacks using spoofed addresses, to prevent illegitimate access to the management infrastructure, and so on.

Ingress filtering can be done, for example, by using access lists or Unicast Reverse Path Forwarding (uRPF). Mechanisms for these are described in [RFC3704].

5.6. Multihoming

Customers may desire multihoming or multi-connecting for a number of reasons [RFC3582].

Mechanisms for multihoming to more than one ISP are still under discussion. One working model would deploy at least one prefix per ISP and choose the prefix from the ISP to which traffic is sent. In addition, tunnels may be used for robustness [RFC3178]. Currently, there are no provider-independent addresses for end-sites. Such addresses would enable IPv4-style multihoming, with associated disadvantages.

Multi-connecting more than once to one ISP is a simple practice, and this can be done, for example, by using BGP with public or private AS numbers and a prefix assigned to the customer.

5.7. Quality of Service

In most networks, quality of service in one form or another is important.

Naturally, the introduction of IPv6 should not impair existing Service Level Agreements (SLAs) or similar quality assurances.

During the deployment of the IPv6 service, the service could be best effort or similar, even if the IPv4 service has an SLA. In the end, both IP versions should be treated equally.

IntServ and DiffServ are equally applicable to IPv6 and IPv4 and work similarly regardless of IP version. Of the two, typically only DiffServ has been implemented.

Many bandwidth provisioning systems operate with IPv4 assumptions, e.g., taking an IPv4 address or (set of) prefixes for which traffic is reserved or preferred. These systems require special attention when introducing IPv6 support in the networks.

6. Network and Service Operation Actions

The network and service operation actions fall into different categories as listed below:

- Set up IPv6 connectivity to upstream providers and peers
- IPv6 network device configuration: for initial configuration and updates
- IPv6 network management
- IPv6 monitoring
- IPv6 customer management
- IPv6 network and service operation security

Some of these items will require an available IPv6 native transport layer and others will not.

As a first step, network device configuration and regular network management operations can be performed over an IPv4 transport, because IPv6 MIBs are also available. Nevertheless, some monitoring functions require the availability of IPv6 transport. This is the case, for instance, when ICMPv6 messages are used by the monitoring applications.

On many platforms, the current inability to retrieve separate IPv4 and IPv6 traffic statistics from dual-stack interfaces for management purposes by using SNMP is an issue.

As a second step, IPv6 transport can be provided for any of these network and service operation facilities.

7. Future Stages

At some point, an ISP may want to change to a service that is IPv6 only, at least in certain parts of its network. This transition creates many new cases into which continued maintenance of the IPv4 service must be factored. Providing an IPv6-only service is not much different from the dual IPv4/IPv6 service described in stage 3 except for the need to phase out the IPv4 service. The delivery of IPv4 services over an IPv6 network and the phaseout of IPv4 are issues

left for a subsequent document. Note that there are some services which will need to maintain IPv4 connectivity (e.g., authoritative and some recursive DNS servers [DNSGUIDE]).

8. Requirements for Follow-On Work

This section tries to summarize the potential items requiring specification in the IETF.

Work items for which an approach was not yet apparent as of this writing are as follows:

- A tunnel server/broker mechanism, for the cases where the customer connection networks cannot be upgraded, needs to be specified [TUNREQS].
- An IPv6 site multihoming mechanism (or multiple ones) needs to be developed.

Work items which were already fast in progress, as of this writing, are as follows:

- 6PE for MPLS was identified as a required mechanism, and this is already in progress [BGPTUNNEL].
- IS-IS for Multiple Topologies was noted as a helpful mechanism in certain environments; however, it is possible to use alternative methods to achieve the same end, so specifying this is not strictly required.

9. Example Networks

This section presents a number of different example networks. These will not necessarily match any existing networks but are intended to be useful even when they do not correspond to specific target networks. The purpose is to exemplify the applicability of the transition mechanisms described in this document to a number of different situations with different prerequisites.

The sample network layout will be the same in each network example. This should be viewed as a specific representation of a generic network with a limited number of network devices. A small number of routers have been used in the examples. However, because the network examples follow the implementation strategies recommended for the generic network scenario, it should be possible to scale the examples to fit a network with an arbitrary number, e.g., several hundreds or thousands of routers.

The routers in the sample network layout are interconnected with each other and with another ISP. The connection to another ISP can be either direct or through an exchange point. A number of customer connection networks are also connected to the routers. Customer connection networks can be, for example, xDSL or cable network equipment.

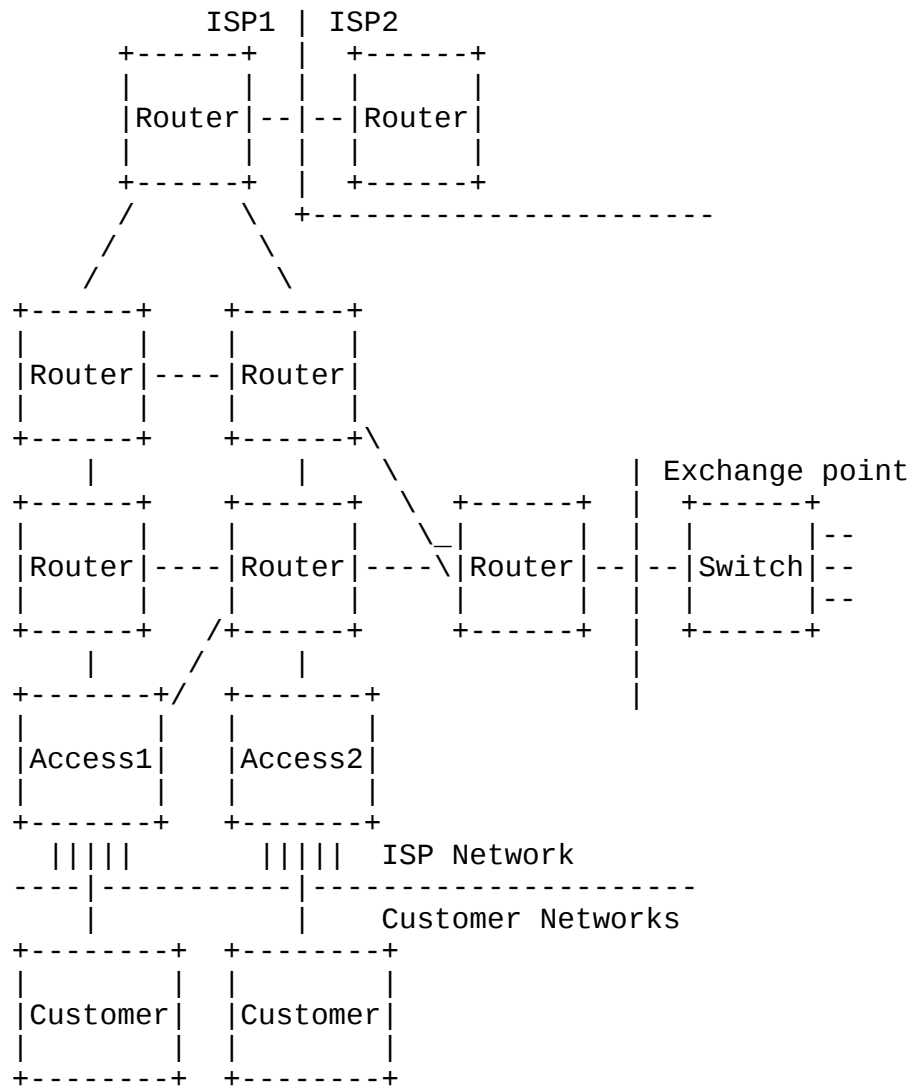


Figure 3: ISP Sample Network Layout

9.1. Example 1

Example 1 presents a network built according to the sample network layout with a native IPv4 backbone. The backbone is running IS-IS and IBGP as routing protocols for internal and external routes, respectively. Multiprotocol BGP is used to exchange routes over the connections to ISP2 and the exchange point. Multicast using PIM-SM routing is present. QoS using DiffServ is deployed.

Access 1 is xDSL connected to the backbone through an access router. The xDSL equipment, except for the access router, is considered to be layer 2 only, e.g., Ethernet or ATM. IPv4 addresses are dynamically assigned to the customer with DHCP. No routing information is exchanged with the customer. Access control and traceability are performed in the access router. Customers are separated into VLANs or separate ATM PVCs up to the access router.

Access 2 is "fiber to the building or home" (FTTB/H) connected directly to the backbone router. This connection is considered layer-3-aware, because it uses layer 3 switches and performs access control and traceability through its layer 3 awareness by using DHCP snooping. IPv4 addresses are dynamically assigned to the customers with DHCP. No routing information is exchanged with the customer.

The actual IPv6 deployment might start by enabling IPv6 on a couple of backbone routers, configuring tunnels between them (if not adjacent) and connecting to a few peers or upstream providers (either through tunnels or at an internet exchange).

After a trial period, the rest of the backbone is upgraded to dual-stack, and IS-IS, without multi-topology extensions (the upgrade order is considered with care), is used as an IPv6 and IPv4 IGP. During an upgrade until IPv6 customers are connected behind a backbone router, the convexity requirement is not critical: The routers will just not be reachable with IPv6. Software supporting IPv6 could be installed even though the routers would not be used for (customer) IPv6 traffic yet. That way, IPv6 could be enabled in the backbone as needed.

Separate IPv6 BGP sessions are built similarly to IPv4. Multicast (through SSM and Embedded-RP) and DiffServ are offered at a later phase of the network, e.g., after a year of stable IPv6 unicast operations.

Offering native service as quickly as possible is important. In the meantime, however, a 6to4 relay may be provided in the meantime for optimized 6to4 connectivity and may also be combined with a tunnel broker for extended functionality. Operating as bridges at Layer 2

only, xDSL equipment does not require changes in CPE: IPv6 connectivity can be offered to the customers by upgrading the PE router to IPv6. In the initial phase, only Router Advertisements are used; DHCPv6 Prefix Delegation can be added as the next step if no other mechanisms are available.

The FTTB/H access has to be upgraded to support access control and traceability in the switches, probably by using DHCP snooping or a similar IPv6 capability, but it also has to be compatible with prefix delegation, not just address assignment. This could, however, lead to the necessity to use DHCPv6 for address assignment.

9.2. Example 2

In example 2, the backbone is running IPv4 with MPLS and is using OSPF and IBGP for internal and external routes, respectively. The connections to ISP2 and the exchange point run BGP to exchange routes. Multicast and QoS are not deployed.

Access 1 is a fixed line, e.g., fiber, connected directly to the backbone. Routing information is in some cases exchanged with CPE at the customer's site; otherwise static routing is used. Access 1 can also be connected to a BGP/MPLS-VPN running in the backbone.

Access 2 is xDSL connected directly to the backbone router. The xDSL is layer 2 only, and access control and traceability are achieved through PPPoE/PPPoA. PPP also provides address assignment. No routing information is exchanged with the customer.

IPv6 deployment might start with an upgrade of a couple of PE routers to support [BGPTUNNEL], as this will allow large-scale IPv6 support without hardware or software upgrades in the core. In a later phase, native IPv6 traffic or IPv6 LSPs would be used in the whole network. In this case, IS-IS or OSPF could be used for the internal routing, and a separate IPv6 BGP session would be run.

For the fixed-line customers, the CPE has to be upgraded, and prefix delegation using DHCPv6 or static assignment would be used. An IPv6 MBGP session would be used when routing information has to be exchanged. In the xDSL case, the same conditions for IP-tunneling apply as in Example 1. In addition to IP-tunneling, a PPP session can be used to offer IPv6 access to a limited number of customers. Later, when clients and servers have been updated, the IPv6 PPP session can be replaced with a combined PPP session for both IPv4 and IPv6. PPP has to be used for address and prefix assignment.

9.3. Example 3

A transit provider offers IP connectivity to other providers, but not to end users or enterprises. IS-IS and IBGP are used internally, and BGP is used externally. Its accesses connect Tier-2 provider cores. No multicast or QoS is used.

As this type of transit provider has a number of customers, who have a large number of customers in turn, it obtains an address allocation from an RIR. The whole backbone can be upgraded to dual-stack in a reasonably short time after a trial with a couple of routers. IPv6 routing is performed by using the same IS-IS process and separate IPv6 BGP sessions.

The ISP provides IPv6 transit to its customers for free, as a competitive advantage. It also provides, at the first phase only, a configured tunnel service with BGP peering to the significant sites and customers (those with an AS number) who are the customers of its customers whenever its own customer networks are not offering IPv6. This is done both to introduce them to IPv6 and to create a beneficial side effect: A bit of extra revenue is generated from its direct customers as the total amount of transited traffic grows.

10. Security Considerations

This document analyzes scenarios and identifies transition mechanisms that could be used for the scenarios. It does not introduce any new security issues. Security considerations of each mechanism are described in the respective documents.

However, a few generic observations are in order.

- o Introducing IPv6 adds new classes of security threats or requires adopting new protocols or operational models than those for IPv4; typically these are generic issues, to be discussed further in other documents, for example, [V6SEC].
- o The more complex the transition mechanisms employed become, the more difficult it will be to manage or analyze their impact on security. Consequently, simple mechanisms are preferable.
- o This document has identified a number of requirements for analysis or further work that should be explicitly considered when adopting IPv6: how to perform access control over shared media or shared ISP customer connection media, how to manage the configuration management security on such environments

(e.g., DHCPv6 authentication keying), and how to manage customer traceability if stateless address autoconfiguration is used.

11. Acknowledgments

This document has greatly benefited from input by Marc Blanchet, Jordi Palet, Francois Le Faucheur, Ronald van der Pol, and Cleve Mickles.

Special thanks to Richard Graveman and Michael Lambert for proofreading the document.

12. Informative References

- [EMBEDRP] Savola, P. and B. Haberman, "Embedding the Rendezvous Point (RP) Address in an IPv6 Multicast Address", RFC 3956, November 2004.
- [MTISIS] Przygienda, T., Naiming Shen, Nischal Sheth, "M-ISIS: Multi Topology (MT) Routing in IS-IS", Work in Progress.
- [RFC2858] Bates, T., Rekhter, Y., Chandra, R., and D. Katz, "Multiprotocol Extensions for BGP-4", RFC 2858, June 2000.
- [RFC2545] Marques, P. and F. Dupont, "Use of BGP-4 Multiprotocol Extensions for IPv6 Inter-Domain Routing", RFC 2545, March 1999.
- [RFC3704] Baker, F. and P. Savola, "Ingress Filtering for Multihomed Networks", BCP 84, RFC 3704, March 2004.
- [RFC3582] Abley, J., Black, B., and V. Gill, "Goals for IPv6 Site-Multihoming Architectures", RFC 3582, August 2003.
- [RFC3178] Hagino, J. and H. Snyder, "IPv6 Multihoming Support at Site Exit Routers", RFC 3178, October 2001.
- [RFC3056] Carpenter, B. and K. Moore, "Connection of IPv6 Domains via IPv4 Clouds", RFC 3056, February 2001.
- [RFC2865] Rigney, C., Willens, S., Rubens, A., and W. Simpson, "Remote Authentication Dial In User Service (RADIUS)", RFC 2865, June 2000.

- [BGPTUNNEL] De Clercq, J., Gastaud, G., Ooms, D., Prevost, S., Le Faucheur, F., "Connecting IPv6 Islands across IPv4 Clouds with BGP", Work in Progress.
- [DUAL-ACCESS] Shirasaki, Y., Miyakawa, S., Yamasaki, T., Takenouchi, A., "A Model of IPv6/IPv4 Dual Stack Internet Access Service", Work in Progress.
- [STEP] Savola, P., "Simple IPv6-in-IPv4 Tunnel Establishment Procedure (STEP)", Work in Progress.
- [TSP] Blanchet, M., "IPv6 Tunnel Broker with Tunnel Setup Protocol (TSP)", Work in Progress.
- [TUNREQS] Palet, J., Nielsen, K., Parent, F., Durand, A., Suryanarayanan, R., and P. Savola, "Goals for Tunneling Configuration", Work in Progress, February 2005.
- [UNMANEVA] Huitema, C., Austein, R., Satapati, S., van der Pol, R., "Evaluation of Transition Mechanisms for Unmanaged Networks", Work in Progress.
- [PROTO41] Palet, J., Olvera, C., Fernandez, D., "Forwarding Protocol 41 in NAT Boxes", Work in Progress.
- [V6SEC] Savola, P., "IPv6 Transition/Co-existence Security Considerations", Work in Progress.
- [DNSGUIDE] Durand, A., Ihren, J., "DNS IPv6 transport operational guidelines", Work in Progress.
- [TEREDO] Huitema, C., "Teredo: Tunneling IPv6 over UDP through NATs", Work in Progress.

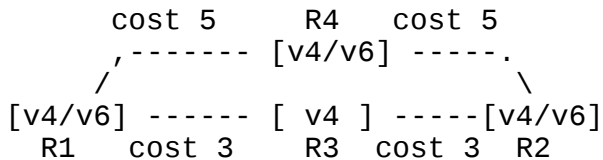
Appendix A: Convexity Requirements in Single Topology IS-IS

The single-topology IS-IS convexity requirements could be summarized, from IPv4/6 perspective, as follows:

- 1) "any IP-independent path from an IPv4 router to any other IPv4 router must only go through routers which are IPv4-capable", and
- 2) "any IP-independent path from an IPv6 router to any other IPv6 router must only go through routers which are IPv6-capable".

As IS-IS is based upon CLNS, these are not trivially accomplished. The single-topology IS-IS builds paths which are agnostic of IP versions.

Consider an example scenario of three IPv4/IPv6-capable routers and an IPv4-only router:



Here the second requirement would not hold. IPv6 packets from R1 to R2 (or vice versa) would go through R3, which does not support IPv6, and the packets would get discarded. By reversing the costs between R1-R3, R3-R2 and R1-R4, R4-R2 the traffic would work in the normal case, but if a link fails and the routing changes to go through R3, the packets would start being discarded again.

Authors' Addresses

Mikael Lind
TeliaSonera
Vitsandsgatan 9B
SE-12386 Farsta, Sweden

Email: mikael.lind@teliasonera.com

Vladimir Ksinant
Thales Communications
160, boulevard de Valmy
92704 Colombes, France

Email: vladimir.ksinant@fr.thalesgroup.com

Soohong Daniel Park
Mobile Platform Laboratory, SAMSUNG Electronics.
416, Maetan-3dong, Paldal-Gu,
Suwon, Gyeonggi-do, Korea

Email: soohong.park@samsung.com

Alain Baudot
France Telecom R&D Division
42, rue des coutures
14066 Caen - FRANCE

Email: alain.baudot@francetelecom.com

Pekka Savola
CSC/FUNET
Espoo, Finland

Email: psavola@funet.fi

Full Copyright Statement

Copyright (C) The Internet Society (2005).

This document is subject to the rights, licenses and restrictions contained in BCP 78, and except as set forth therein, the authors retain all their rights.

This document and the information contained herein are provided on an "AS IS" basis and THE CONTRIBUTOR, THE ORGANIZATION HE/SHE REPRESENTS OR IS SPONSORED BY (IF ANY), THE INTERNET SOCIETY AND THE INTERNET ENGINEERING TASK FORCE DISCLAIM ALL WARRANTIES, EXPRESS OR IMPLIED, INCLUDING BUT NOT LIMITED TO ANY WARRANTY THAT THE USE OF THE INFORMATION HEREIN WILL NOT INFRINGE ANY RIGHTS OR ANY IMPLIED WARRANTIES OF MERCHANTABILITY OR FITNESS FOR A PARTICULAR PURPOSE.

Intellectual Property

The IETF takes no position regarding the validity or scope of any Intellectual Property Rights or other rights that might be claimed to pertain to the implementation or use of the technology described in this document or the extent to which any license under such rights might or might not be available; nor does it represent that it has made any independent effort to identify any such rights. Information on the procedures with respect to rights in RFC documents can be found in BCP 78 and BCP 79.

Copies of IPR disclosures made to the IETF Secretariat and any assurances of licenses to be made available, or the result of an attempt made to obtain a general license or permission for the use of such proprietary rights by implementers or users of this specification can be obtained from the IETF on-line IPR repository at <http://www.ietf.org/ipr>.

The IETF invites any interested party to bring to its attention any copyrights, patents or patent applications, or other proprietary rights that may cover technology that may be required to implement this standard. Please address the information to the IETF at ietf-ipr@ietf.org.

Acknowledgement

Funding for the RFC Editor function is currently provided by the Internet Society.

