

Network Working Group
Request for Comments: 2764
Category: Informational

B. Gleeson
A. Lin
Nortel Networks
J. Heinanen
Telia Finland
G. Armitage
A. Malis
Lucent Technologies
February 2000

A Framework for IP Based Virtual Private 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 (2000). All Rights Reserved.

IESG Note

This document is not the product of an IETF Working Group. The IETF currently has no effort underway to standardize a specific VPN framework.

Abstract

This document describes a framework for Virtual Private Networks (VPNs) running across IP backbones. It discusses the various different types of VPNs, their respective requirements, and proposes specific mechanisms that could be used to implement each type of VPN using existing or proposed specifications. The objective of this document is to serve as a framework for related protocol development in order to develop the full set of specifications required for widespread deployment of interoperable VPN solutions.

Table of Contents

1.0 Introduction	4
2.0 VPN Application and Implementation Requirements	5
2.1 General VPN Requirements	5
2.1.1 Opaque Packet Transport:	6
2.1.2 Data Security	7
2.1.3 Quality of Service Guarantees	7
2.1.4 Tunneling Mechanism	8
2.2 CPE and Network Based VPNs	8
2.3 VPNs and Extranets	9
3.0 VPN Tunneling	10
3.1 Tunneling Protocol Requirements for VPNs	11
3.1.1 Multiplexing	11
3.1.2 Signalling Protocol	12
3.1.3 Data Security	13
3.1.4 Multiprotocol Transport	14
3.1.5 Frame Sequencing	14
3.1.6 Tunnel Maintenance	15
3.1.7 Large MTUs	16
3.1.8 Minimization of Tunnel Overhead	16
3.1.9 Flow and congestion control	17
3.1.10 QoS / Traffic Management	17
3.2 Recommendations	18
4.0 VPN Types: Virtual Leased Lines	18
5.0 VPN Types: Virtual Private Routed Networks	20
5.1 VPRN Characteristics	20
5.1.1 Topology	23
5.1.2 Addressing	24
5.1.3 Forwarding	24
5.1.4 Multiple concurrent VPRN connectivity	24
5.2 VPRN Related Work	24
5.3 VPRN Generic Requirements	25
5.3.1 VPN Identifier	26
5.3.2 VPN Membership Information Configuration	27
5.3.2.1 Directory Lookup	27
5.3.2.2 Explicit Management Configuration	28
5.3.2.3 Piggybacking in Routing Protocols	28
5.3.3 Stub Link Reachability Information	30
5.3.3.1 Stub Link Connectivity Scenarios	30
5.3.3.1.1 Dual VPRN and Internet Connectivity	30
5.3.3.1.2 VPRN Connectivity Only	30
5.3.3.1.3 Multihomed Connectivity	31
5.3.3.1.4 Backdoor Links	31
5.3.3.1 Routing Protocol Instance	31
5.3.3.2 Configuration	33
5.3.3.3 ISP Administered Addresses	33
5.3.3.4 MPLS Label Distribution Protocol	33

5.3.4 Intra-VPN Reachability Information	34
5.3.4.1 Directory Lookup	34
5.3.4.2 Explicit Configuration	34
5.3.4.3 Local Intra-VPN Routing Instantiations	34
5.3.4.4 Link Reachability Protocol	35
5.3.4.5 Piggybacking in IP Backbone Routing Protocols	36
5.3.5 Tunneling Mechanisms	36
5.4 Multihomed Stub Routers	37
5.5 Multicast Support	38
5.5.1 Edge Replication	38
5.5.2 Native Multicast Support	39
5.6 Recommendations	40
6.0 VPN Types: Virtual Private Dial Networks	41
6.1 L2TP protocol characteristics	41
6.1.1 Multiplexing	41
6.1.2 Signalling	42
6.1.3 Data Security	42
6.1.4 Multiprotocol Transport	42
6.1.5 Sequencing	42
6.1.6 Tunnel Maintenance	43
6.1.7 Large MTUs	43
6.1.8 Tunnel Overhead	43
6.1.9 Flow and Congestion Control	43
6.1.10 QoS / Traffic Management	43
6.1.11 Miscellaneous	44
6.2 Compulsory Tunneling	44
6.3 Voluntary Tunnels	46
6.3.1 Issues with Use of L2TP for Voluntary Tunnels	46
6.3.2 Issues with Use of IPsec for Voluntary Tunnels	48
6.4 Networked Host Support	49
6.4.1 Extension of PPP to Hosts Through L2TP	49
6.4.2 Extension of PPP Directly to Hosts:	49
6.4.3 Use of IPsec	50
6.5 Recommendations	50
7.0 VPN Types: Virtual Private LAN Segment	50
7.1 VPLS Requirements	51
7.1.1 Tunneling Protocols	51
7.1.2 Multicast and Broadcast Support	52
7.1.3 VPLS Membership Configuration and Topology	52
7.1.4 CPE Stub Node Types	52
7.1.5 Stub Link Packet Encapsulation	53
7.1.5.1 Bridge CPE	53
7.1.5.2 Router CPE	53
7.1.6 CPE Addressing and Address Resolution	53
7.1.6.1 Bridge CPE	53
7.1.6.2 Router CPE	54
7.1.7 VPLS Edge Node Forwarding and Reachability Mechanisms	54
7.1.7.1 Bridge CPE	54

7.1.7.2 Router CPE	54
7.2 Recommendations	55
8.0 Summary of Recommendations	55
9.0 Security Considerations	56
10.0 Acknowledgements	56
11.0 References	56
12.0 Author Information	61
13.0 Full Copyright Statement	62

1.0 Introduction

This document describes a framework for Virtual Private Networks (VPNs) running across IP backbones. It discusses the various different types of VPNs, their respective requirements, and proposes specific mechanisms that could be used to implement each type of VPN using existing or proposed specifications. The objective of this document is to serve as a framework for related protocol development in order to develop the full set of specifications required for widespread deployment of interoperable VPN solutions.

There is currently significant interest in the deployment of virtual private networks across IP backbone facilities. The widespread deployment of VPNs has been hampered, however, by the lack of interoperable implementations, which, in turn, derives from the lack of general agreement on the definition and scope of VPNs and confusion over the wide variety of solutions that are all described by the term VPN. In the context of this document, a VPN is simply defined as the 'emulation of a private Wide Area Network (WAN) facility using IP facilities' (including the public Internet, or private IP backbones). As such, there are as many types of VPNs as there are types of WANs, hence the confusion over what exactly constitutes a VPN.

In this document a VPN is modeled as a connectivity object. Hosts may be attached to a VPN, and VPNs may be interconnected together, in the same manner as hosts today attach to physical networks, and physical networks are interconnected together (e.g., via bridges or routers). Many aspects of networking, such as addressing, forwarding mechanism, learning and advertising reachability, quality of service (QoS), security, and firewalling, have common solutions across both physical and virtual networks, and many issues that arise in the discussion of VPNs have direct analogues with those issues as implemented in physical networks. The introduction of VPNs does not create the need to reinvent networking, or to introduce entirely new paradigms that have no direct analogue with existing physical networks. Instead it is often useful to first examine how a particular issue is handled in a physical network environment, and then apply the same principle to an environment which contains

virtual as well as physical networks, and to develop appropriate extensions and enhancements when necessary. Clearly having mechanisms that are common across both physical and virtual networks facilitates the introduction of VPNs into existing networks, and also reduces the effort needed for both standards and product development, since existing solutions can be leveraged.

This framework document proposes a taxonomy of a specific set of VPN types, showing the specific applications of each, their specific requirements, and the specific types of mechanisms that may be most appropriate for their implementation. The intent of this document is to serve as a framework to guide a coherent discussion of the specific modifications that may be needed to existing IP mechanisms in order to develop a full range of interoperable VPN solutions.

The document first discusses the likely expectations customers have of any type of VPN, and the implications of these for the ways in which VPNs can be implemented. It also discusses the distinctions between Customer Premises Equipment (CPE) based solutions, and network based solutions. Thereafter it presents a taxonomy of the various VPN types and their respective requirements. It also outlines suggested approaches to their implementation, hence also pointing to areas for future standardization.

Note also that this document only discusses implementations of VPNs across IP backbones, be they private IP networks, or the public Internet. The models and mechanisms described here are intended to apply to both IPV4 and IPV6 backbones. This document specifically does not discuss means of constructing VPNs using native mappings onto switched backbones - e.g., VPNs constructed using the LAN Emulation over ATM (LANE) [1] or Multiprotocol over ATM (MPOA) [2] protocols operating over ATM backbones. Where IP backbones are constructed using such protocols, by interconnecting routers over the switched backbone, the VPNs discussed operate on top of this IP network, and hence do not directly utilize the native mechanisms of the underlying backbone. Native VPNs are restricted to the scope of the underlying backbone, whereas IP based VPNs can extend to the extent of IP reachability. Native VPN protocols are clearly outside the scope of the IETF, and may be tackled by such bodies as the ATM Forum.

2.0 VPN Application and Implementation Requirements

2.1 General VPN Requirements

There is growing interest in the use of IP VPNs as a more cost effective means of building and deploying private communication networks for multi-site communication than with existing approaches.

Existing private networks can be generally categorized into two types - dedicated WANs that permanently connect together multiple sites, and dial networks, that allow on-demand connections through the Public Switched Telephone Network (PSTN) to one or more sites in the private network.

WANs are typically implemented using leased lines or dedicated circuits - for instance, Frame Relay or ATM connections - between the multiple sites. CPE routers or switches at the various sites connect these dedicated facilities together and allow for connectivity across the network. Given the cost and complexity of such dedicated facilities and the complexity of CPE device configuration, such networks are generally not fully meshed, but instead have some form of hierarchical topology. For example remote offices could be connected directly to the nearest regional office, with the regional offices connected together in some form of full or partial mesh.

Private dial networks are used to allow remote users to connect into an enterprise network using PSTN or Integrated Services Digital Network (ISDN) links. Typically, this is done through the deployment of Network Access Servers (NASs) at one or more central sites. Users dial into such NASs, which interact with Authentication, Authorization, and Accounting (AAA) servers to verify the identity of the user, and the set of services that the user is authorized to receive.

In recent times, as more businesses have found the need for high speed Internet connections to their private corporate networks, there has been significant interest in the deployment of CPE based VPNs running across the Internet. This has been driven typically by the ubiquity and distance insensitive pricing of current Internet services, that can result in significantly lower costs than typical dedicated or leased line services.

The notion of using the Internet for private communications is not new, and many techniques, such as controlled route leaking, have been used for this purpose [3]. Only in recent times, however, have the appropriate IP mechanisms needed to meet customer requirements for VPNs all come together. These requirements include the following:

2.1.1 Opaque Packet Transport:

The traffic carried within a VPN may have no relation to the traffic on the IP backbone, either because the traffic is multiprotocol, or because the customer's IP network may use IP addressing unrelated to that of the IP backbone on which the traffic is transported. In particular, the customer's IP network may use non-unique, private IP addressing [4].

2.1.2 Data Security

In general customers using VPNs require some form of data security. There are different trust models applicable to the use of VPNs. One such model is where the customer does not trust the service provider to provide any form of security, and instead implements a VPN using CPE devices that implement firewall functionality and that are connected together using secure tunnels. In this case the service provider is used solely for IP packet transport.

An alternative model is where the customer trusts the service provider to provide a secure managed VPN service. This is similar to the trust involved when a customer utilizes a public switched Frame Relay or ATM service, in that the customer trusts that packets will not be misdirected, injected into the network in an unauthorized manner, snooped on, modified in transit, or subjected to traffic analysis by unauthorized parties.

With this model providing firewall functionality and secure packet transport services is the responsibility of the service provider. Different levels of security may be needed within the provider backbone, depending on the deployment scenario used. If the VPN traffic is contained within a single provider's IP backbone then strong security mechanisms, such as those provided by the IP Security protocol suite (IPSec) [5], may not be necessary for tunnels between backbone nodes. If the VPN traffic traverses networks or equipment owned by multiple administrations then strong security mechanisms may be appropriate. Also a strong level of security may be applied by a provider to customer traffic to address a customer perception that IP networks, and particularly the Internet, are insecure. Whether or not this perception is correct it is one that must be addressed by the VPN implementation.

2.1.3 Quality of Service Guarantees

In addition to ensuring communication privacy, existing private networking techniques, building upon physical or link layer mechanisms, also offer various types of quality of service guarantees. In particular, leased and dial up lines offer both bandwidth and latency guarantees, while dedicated connection technologies like ATM and Frame Relay have extensive mechanisms for similar guarantees. As IP based VPNs become more widely deployed, there will be market demand for similar guarantees, in order to ensure end to end application transparency. While the ability of IP based VPNs to offer such guarantees will depend greatly upon the commensurate capabilities of the underlying IP backbones, a VPN framework must also address the means by which VPN systems can utilize such capabilities, as they evolve.

2.1.4 Tunneling Mechanism

Together, the first two of the requirements listed above imply that VPNs must be implemented through some form of IP tunneling mechanism, where the packet formats and/or the addressing used within the VPN can be unrelated to that used to route the tunneled packets across the IP backbone. Such tunnels, depending upon their form, can provide some level of intrinsic data security, or this can also be enhanced using other mechanisms (e.g., IPSec).

Furthermore, as discussed later, such tunneling mechanisms can also be mapped into evolving IP traffic management mechanisms. There are already defined a large number of IP tunneling mechanisms. Some of these are well suited to VPN applications, as discussed in section 3.0.

2.2 CPE and Network Based VPNs

Most current VPN implementations are based on CPE equipment. VPN capabilities are being integrated into a wide variety of CPE devices, ranging from firewalls to WAN edge routers and specialized VPN termination devices. Such equipment may be bought and deployed by customers, or may be deployed (and often remotely managed) by service providers in an outsourcing service.

There is also significant interest in 'network based VPNs', where the operation of the VPN is outsourced to an Internet Service Provider (ISP), and is implemented on network as opposed to CPE equipment. There is significant interest in such solutions both by customers seeking to reduce support costs and by ISPs seeking new revenue sources. Supporting VPNs in the network allows the use of particular mechanisms which may lead to highly efficient and cost effective VPN solutions, with common equipment and operations support amortized across large numbers of customers.

Most of the mechanisms discussed below can apply to either CPE based or network based VPNs. However particular mechanisms are likely to prove applicable only to the latter, since they leverage tools (e.g., piggybacking on routing protocols) which are accessible only to ISPs and which are unlikely to be made available to any customer, or even hosted on ISP owned and operated CPE, due to the problems of coordinating joint management of the CPE gear by both the ISP and the customer. This document will indicate which techniques are likely to apply only to network based VPNs.

2.3 VPNs and Extranets

The term 'extranet' is commonly used to refer to a scenario whereby two or more companies have networked access to a limited amount of each other's corporate data. For example a manufacturing company might use an extranet for its suppliers to allow it to query databases for the pricing and availability of components, and then to order and track the status of outstanding orders. Another example is joint software development, for instance, company A allows one development group within company B to access its operating system source code, and company B allows one development group in company A to access its security software. Note that the access policies can get arbitrarily complex. For example company B may internally restrict access to its security software to groups in certain geographic locations to comply with export control laws, for example.

A key feature of an extranet is thus the control of who can access what data, and this is essentially a policy decision. Policy decisions are typically enforced today at the interconnection points between different domains, for example between a private network and the Internet, or between a software test lab and the rest of the company network. The enforcement may be done via a firewall, router with access list functionality, application gateway, or any similar device capable of applying policy to transit traffic. Policy controls may be implemented within a corporate network, in addition to between corporate networks. Also the interconnections between networks could be a set of bilateral links, or could be a separate network, perhaps maintained by an industry consortium. This separate network could itself be a VPN or a physical network.

Introducing VPNs into a network does not require any change to this model. Policy can be enforced between two VPNs, or between a VPN and the Internet, in exactly the same manner as is done today without VPNs. For example two VPNs could be interconnected, which each administration locally imposing its own policy controls, via a firewall, on all traffic that enters its VPN from the outside, whether from another VPN or from the Internet.

This model of a VPN provides for a separation of policy from the underlying mode of packet transport used. For example, a router may direct voice traffic to ATM Virtual Channel Connections (VCCs) for guaranteed QoS, non-local internal company traffic to secure tunnels, and other traffic to a link to the Internet. In the past the secure tunnels may have been frame relay circuits, now they may also be secure IP tunnels or MPLS Label Switched Paths (LSPs)

Other models of a VPN are also possible. For example there is a model whereby a set of application flows is mapped into a VPN. As the policy rules imposed by a network administrator can get quite complex, the number of distinct sets of application flows that are used in the policy rulebase, and hence the number of VPNs, can thus grow quite large, and there can be multiple overlapping VPNs. However there is little to be gained by introducing such new complexity into a network. Instead a VPN should be viewed as a direct analogue to a physical network, as this allows the leveraging of existing protocols and procedures, and the current expertise and skill sets of network administrators and customers.

3.0 VPN Tunneling

As noted above in section 2.1, VPNs must be implemented using some form of tunneling mechanism. This section looks at the generic requirements for such VPN tunneling mechanisms. A number of characteristics and aspects common to any link layer protocol are taken and compared with the features offered by existing tunneling protocols. This provides a basis for comparing different protocols and is also useful to highlight areas where existing tunneling protocols could benefit from extensions to better support their operation in a VPN environment.

An IP tunnel connecting two VPN endpoints is a basic building block from which a variety of different VPN services can be constructed. An IP tunnel operates as an overlay across the IP backbone, and the traffic sent through the tunnel is opaque to the underlying IP backbone. In effect the IP backbone is being used as a link layer technology, and the tunnel forms a point-to-point link.

A VPN device may terminate multiple IP tunnels and forward packets between these tunnels and other network interfaces in different ways. In the discussion of different types of VPNs, in later sections of this document, the primary distinguishing characteristic of these different types is the manner in which packets are forwarded between interfaces (e.g., bridged or routed). There is a direct analogy with how existing networking devices are characterized today. A two-port repeater just forwards packets between its ports, and does not examine the contents of the packet. A bridge forwards packets using Media Access Control (MAC) layer information contained in the packet, while a router forwards packets using layer 3 addressing information contained in the packet. Each of these three scenarios has a direct VPN analogue, as discussed later. Note that an IP tunnel is viewed as just another sort of link, which can be concatenated with another link, bound to a bridge forwarding table, or bound to an IP forwarding table, depending on the type of VPN.

The following sections look at the requirements for a generic IP tunneling protocol that can be used as a basic building block to construct different types of VPNs.

3.1 Tunneling Protocol Requirements for VPNs

There are numerous IP tunneling mechanisms, including IP/IP [6], Generic Routing Encapsulation (GRE) tunnels [7], Layer 2 Tunneling Protocol (L2TP) [8], IPSec [5], and Multiprotocol Label Switching (MPLS) [9]. Note that while some of these protocols are not often thought of as tunneling protocols, they do each allow for opaque transport of frames as packet payload across an IP network, with forwarding disjoint from the address fields of the encapsulated packets.

Note, however, that there is one significant distinction between each of the IP tunneling protocols mentioned above, and MPLS. MPLS can be viewed as a specific link layer for IP, insofar as MPLS specific mechanisms apply only within the scope of an MPLS network, whereas IP based mechanisms extend to the extent of IP reachability. As such, VPN mechanisms built directly upon MPLS tunneling mechanisms cannot, by definition, extend outside the scope of MPLS networks, any more so than, for instance, ATM based mechanisms such as LANE can extend outside of ATM networks. Note however, that an MPLS network can span many different link layer technologies, and so, like an IP network, its scope is not limited by the specific link layers used. A number of proposals for defining a set of mechanisms to allow for interoperable VPNs specifically over MPLS networks have also been produced ([10] [11] [12] [13], [14] and [15]).

There are a number of desirable requirements for a VPN tunneling mechanism, however, that are not all met by the existing tunneling mechanisms. These requirements include:

3.1.1 Multiplexing

There are cases where multiple VPN tunnels may be needed between the same two IP endpoints. This may be needed, for instance, in cases where the VPNs are network based, and each end point supports multiple customers. Traffic for different customers travels over separate tunnels between the same two physical devices. A multiplexing field is needed to distinguish which packets belong to which tunnel. Sharing a tunnel in this manner may also reduce the latency and processing burden of tunnel set up. Of the existing IP tunneling mechanisms, L2TP (via the tunnel-id and session-id fields), MPLS (via the label) and IPSec (via the Security Parameter Index (SPI) field) have a multiplexing mechanism. Strictly speaking GRE does not have a multiplexing field. However the key field, which was

intended to be used for authenticating the source of a packet, has sometimes been used as a multiplexing field. IP/IP does not have a multiplexing field.

The IETF [16] and the ATM Forum [17] have standardized on a single format for a globally unique identifier used to identify a VPN (a VPN-ID). A VPN-ID can be used in the control plane, to bind a tunnel to a VPN at tunnel establishment time, or in the data plane, to identify the VPN associated with a packet, on a per-packet basis. In the data plane a VPN encapsulation header can be used by MPLS, MPOA and other tunneling mechanisms to aggregate packets for different VPNs over a single tunnel. In this case an explicit indication of VPN-ID is included with every packet, and no use is made of any tunnel specific multiplexing field. In the control plane a VPN-ID field can be included in any tunnel establishment signalling protocol to allow for the association of a tunnel (e.g., as identified by the SPI field) with a VPN. In this case there is no need for a VPN-ID to be included with every data packet. This is discussed further in section 5.3.1.

3.1.2 Signalling Protocol

There is some configuration information that must be known by an end point in advance of tunnel establishment, such as the IP address of the remote end point, and any relevant tunnel attributes required, such as the level of security needed. Once this information is available, the actual tunnel establishment can be completed in one of two ways - via a management operation, or via a signalling protocol that allows tunnels to be established dynamically.

An example of a management operation would be to use an SNMP Management Information Base (MIB) to configure various tunneling parameters, e.g., MPLS labels, source addresses to use for IP/IP or GRE tunnels, L2TP tunnel-ids and session-ids, or security association parameters for IPsec.

Using a signalling protocol can significantly reduce the management burden however, and as such, is essential in many deployment scenarios. It reduces the amount of configuration needed, and also reduces the management co-ordination needed if a VPN spans multiple administrative domains. For example, the value of the multiplexing field, described above, is local to the node assigning the value, and can be kept local if distributed via a signalling protocol, rather than being first configured into a management station and then distributed to the relevant nodes. A signalling protocol also allows nodes that are mobile or are only intermittently connected to establish tunnels on demand.

When used in a VPN environment a signalling protocol should allow for the transport of a VPN-ID to allow the resulting tunnel to be associated with a particular VPN. It should also allow tunnel attributes to be exchanged or negotiated, for example the use of frame sequencing or the use of multiprotocol transport. Note that the role of the signalling protocol need only be to negotiate tunnel attributes, not to carry information about how the tunnel is used, for example whether the frames carried in the tunnel are to be forwarded at layer 2 or layer 3. (This is similar to Q.2931 ATM signalling - the same signalling protocol is used to set up Classical IP logical subnetworks as well as for LANE emulated LANs.

Of the various IP tunneling protocols, the following ones support a signalling protocol that could be adapted for this purpose: L2TP (the L2TP control protocol), IPsec (the Internet Key Exchange (IKE) protocol [18]), and GRE (as used with mobile-ip tunneling [19]). Also there are two MPLS signalling protocols that can be used to establish LSP tunnels. One uses extensions to the MPLS Label Distribution Protocol (LDP) protocol [20], called Constraint-Based Routing LDP (CR-LDP) [21], and the other uses extensions to the Resource Reservation Protocol (RSVP) for LSP tunnels [22].

3.1.3 Data Security

A VPN tunneling protocol must support mechanisms to allow for whatever level of security may be desired by customers, including authentication and/or encryption of various strengths. None of the tunneling mechanisms discussed, other than IPsec, have intrinsic security mechanisms, but rely upon the security characteristics of the underlying IP backbone. In particular, MPLS relies upon the explicit labeling of label switched paths to ensure that packets cannot be misdirected, while the other tunneling mechanisms can all be secured through the use of IPsec. For VPNs implemented over non-IP backbones (e.g., MPOA, Frame Relay or ATM virtual circuits), data security is implicitly provided by the layer two switch infrastructure.

Overall VPN security is not just a capability of the tunnels alone, but has to be viewed in the broader context of how packets are forwarded onto those tunnels. For example with VPRNs implemented with virtual routers, the use of separate routing and forwarding table instances ensures the isolation of traffic between VPNs. Packets on one VPN cannot be misrouted to a tunnel on a second VPN since those tunnels are not visible to the forwarding table of the first VPN.

If some form of signalling mechanism is used by one VPN end point to dynamically establish a tunnel with another endpoint, then there is a requirement to be able to authenticate the party attempting the tunnel establishment. IPSec has an array of schemes for this purpose, allowing, for example, authentication to be based on pre-shared keys, or to use digital signatures and certificates. Other tunneling schemes have weaker forms of authentication. In some cases no authentication may be needed, for example if the tunnels are provisioned, rather than dynamically established, or if the trust model in use does not require it.

Currently the IPSec Encapsulating Security Payload (ESP) protocol [23] can be used to establish SAs that support either encryption or authentication or both. However the protocol specification precludes the use of an SA where neither encryption or authentication is used. In a VPN environment this "null/null" option is useful, since other aspects of the protocol (e.g., that it supports tunneling and multiplexing) may be all that is required. In effect the "null/null" option can be viewed as just another level of data security.

3.1.4 Multiprotocol Transport

In many applications of VPNs, the VPN may carry opaque, multiprotocol traffic. As such, the tunneling protocol used must also support multiprotocol transport. L2TP is designed to transport Point-to-Point Protocol (PPP) [24] packets, and thus can be used to carry multiprotocol traffic since PPP itself is multiprotocol. GRE also provides for the identification of the protocol being tunneled. IP/IP and IPSec tunnels have no such protocol identification field, since the traffic being tunneled is assumed to be IP.

It is possible to extend the IPSec protocol suite to allow for the transport of multiprotocol packets. This can be achieved, for example, by extending the signalling component of IPSec - IKE, to indicate the protocol type of the traffic being tunneled, or to carry a packet multiplexing header (e.g., an LLC/SNAP header or GRE header) with each tunneled packet. This approach is similar to that used for the same purpose in ATM networks, where signalling is used to indicate the encapsulation used on the VCC, and where packets sent on the VCC can use either an LLC/SNAP header or be placed directly into the AAL5 payload, the latter being known as VC-multiplexing (see [25]).

3.1.5 Frame Sequencing

One quality of service attribute required by customers of a VPN may be frame sequencing, matching the equivalent characteristic of physical leased lines or dedicated connections. Sequencing may be

required for the efficient operation of particular end to end protocols or applications. In order to implement frame sequencing, the tunneling mechanism must support a sequencing field. Both L2TP and GRE have such a field. IPSec has a sequence number field, but it is used by a receiver to perform an anti-replay check, not to guarantee in-order delivery of packets.

It is possible to extend IPSec to allow the use of the existing sequence field to guarantee in-order delivery of packets. This can be achieved, for example, by using IKE to negotiate whether or not sequencing is to be used, and to define an end point behaviour which preserves packet sequencing.

3.1.6 Tunnel Maintenance

The VPN end points must monitor the operation of the VPN tunnels to ensure that connectivity has not been lost, and to take appropriate action (such as route recalculation) if there has been a failure.

There are two approaches possible. One is for the tunneling protocol itself to periodically check in-band for loss of connectivity, and to provide an explicit indication of failure. For example L2TP has an optional keep-alive mechanism to detect non-operational tunnels.

The other approach does not require the tunneling protocol itself to perform this function, but relies on the operation of some out-of-band mechanism to determine loss of connectivity. For example if a routing protocol such as Routing Information Protocol (RIP) [26] or Open Shortest Path First (OSPF) [27] is run over a tunnel mesh, a failure to hear from a neighbor within a certain period of time will result in the routing protocol declaring the tunnel to be down. Another out-of-band approach is to perform regular ICMP pings with a peer. This is generally sufficient assurance that the tunnel is operational, due to the fact the tunnel also runs across the same IP backbone.

When tunnels are established dynamically a distinction needs to be drawn between the static and dynamic tunnel information needed. Before a tunnel can be established some static information is needed by a node, such as the identify of the remote end point and the attributes of the tunnel to propose and accept. This is typically put in place as a result of a configuration operation. As a result of the signalling exchange to establish a tunnel, some dynamic state is established in each end point, such as the value of the multiplexing field or keys to be used. For example with IPSec, the establishment of a Security Association (SA) puts in place the keys to be used for the lifetime of that SA.

Different policies may be used as to when to trigger the establishment of a dynamic tunnel. One approach is to use a data-driven approach and to trigger tunnel establishment whenever there is data to be transferred, and to timeout the tunnel due to inactivity. This approach is particularly useful if resources for the tunnel are being allocated in the network for QoS purposes. Another approach is to trigger tunnel establishment whenever the static tunnel configuration information is installed, and to attempt to keep the tunnel up all the time.

3.1.7 Large MTUs

An IP tunnel has an associated Maximum Transmission Unit (MTU), just like a regular link. It is conceivable that this MTU may be larger than the MTU of one or more individual hops along the path between tunnel endpoints. If so, some form of frame fragmentation will be required within the tunnel.

If the frame to be transferred is mapped into one IP datagram, normal IP fragmentation will occur when the IP datagram reaches a hop with an MTU smaller than the IP tunnel's MTU. This can have undesirable performance implications at the router performing such mid-tunnel fragmentation.

An alternative approach is for the tunneling protocol itself to incorporate a segmentation and reassembly capability that operates at the tunnel level, perhaps using the tunnel sequence number and an end-of-message marker of some sort. (Note that multilink PPP uses a mechanism similar to this to fragment packets). This avoids IP level fragmentation within the tunnel itself. None of the existing tunneling protocols support such a mechanism.

3.1.8 Minimization of Tunnel Overhead

There is clearly benefit in minimizing the overhead of any tunneling mechanisms. This is particularly important for the transport of jitter and latency sensitive traffic such as packetized voice and video. On the other hand, the use of security mechanisms, such as IPSec, do impose their own overhead, hence the objective should be to minimize overhead over and above that needed for security, and to not burden those tunnels in which security is not mandatory with unnecessary overhead.

One area where the amount of overhead may be significant is when voluntary tunneling is used for dial-up remote clients connecting to a VPN, due to the typically low bandwidth of dial-up links. This is discussed further in section 6.3.

3.1.9 Flow and congestion control

During the development of the L2TP protocol procedures were developed for flow and congestion control. These were necessitated primarily because of the need to provide adequate performance over lossy networks when PPP compression is used, which, unlike IP Payload Compression Protocol (IPComp) [28], is stateful across packets. Another motivation was to accommodate devices with very little buffering, used for example to terminate low speed dial-up lines. However the flow and congestion control mechanisms defined in the final version of the L2TP specification are used only for the control channels, and not for data traffic.

In general the interactions between multiple layers of flow and congestion control schemes can be very complex. Given the predominance of TCP traffic in today's networks and the fact that TCP has its own end-to-end flow and congestion control mechanisms, it is not clear that there is much benefit to implementing similar mechanisms within tunneling protocols. Good flow and congestion control schemes, that can adapt to a wide variety of network conditions and deployment scenarios are complex to develop and test, both in themselves and in understanding the interaction with other schemes that may be running in parallel. There may be some benefit, however, in having the capability whereby a sender can shape traffic to the capacity of a receiver in some manner, and in providing the protocol mechanisms to allow a receiver to signal its capabilities to a sender. This is an area that may benefit from further study.

Note also the work of the Performance Implications of Link Characteristics (PILC) working group of the IETF, which is examining how the properties of different network links can have an impact on the performance of Internet protocols operating over those links.

3.1.10 QoS / Traffic Management

As noted above, customers may require that VPNs yield similar behaviour to physical leased lines or dedicated connections with respect to such QoS parameters as loss rates, jitter, latency and bandwidth guarantees. How such guarantees could be delivered will, in general, be a function of the traffic management characteristics of the VPN nodes themselves, and the access and backbone networks across which they are connected.

A full discussion of QoS and VPNs is outside the scope of this document, however by modeling a VPN tunnel as just another type of link layer, many of the existing mechanisms developed for ensuring QoS over physical links can also be applied. For example at a VPN node, the mechanisms of policing, marking, queuing, shaping and

scheduling can all be applied to VPN traffic with VPN-specific parameters, queues and interfaces, just as for non-VPN traffic. The techniques developed for Diffserv, Intserv and for traffic engineering in MPLS are also applicable. See also [29] for a discussion of QoS and VPNs.

It should be noted, however, that this model of tunnel operation is not necessarily consistent with the way in which specific tunneling protocols are currently modeled. While a model is an aid to comprehension, and not part of a protocol specification, having differing models can complicate discussions, particularly if a model is misinterpreted as being part of a protocol specification or as constraining choice of implementation method. For example, IPsec tunnel processing can be modeled both as an interface and as an attribute of a particular packet flow.

3.2 Recommendations

IPsec is needed whenever there is a requirement for strong encryption or strong authentication. It also supports multiplexing and a signalling protocol - IKE. However extending the IPsec protocol suite to also cover the following areas would be beneficial, in order to better support the tunneling requirements of a VPN environment.

- the transport of a VPN-ID when establishing an SA (3.1.2)
- a null encryption and null authentication option (3.1.3)
- multiprotocol operation (3.1.4)
- frame sequencing (3.1.5)

L2TP provides no data security by itself, and any PPP security mechanisms used do not apply to the L2TP protocol itself, so that in order for strong security to be provided L2TP must run over IPsec. Defining specific modes of operation for IPsec when it is used to support L2TP traffic will aid interoperability. This is currently a work item for the proposed L2TP working group.

4.0 VPN Types: Virtual Leased Lines

The simplest form of a VPN is a 'Virtual Leased Line' (VLL) service. In this case a point-to-point link is provided to a customer, connecting two CPE devices, as illustrated below. The link layer type used to connect the CPE devices to the ISP nodes can be any link layer type, for example an ATM VCC or a Frame Relay circuit. The CPE devices can be either routers bridges or hosts.

The two ISP nodes are both connected to an IP network, and an IP tunnel is set up between them. Each ISP node is configured to bind the stub link and the IP tunnel together at layer 2 (e.g., an ATM VCC and the IP tunnel). Frames are relayed between the two links. For example the ATM Adaptation Layer 5 (AAL5) payload is taken and encapsulated in an IPsec tunnel, and vice versa. The contents of the AAL5 payload are opaque to the ISP node, and are not examined there.

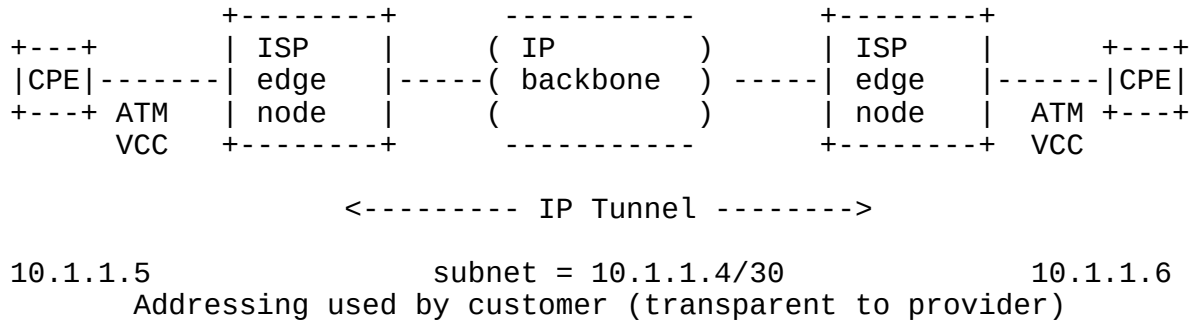


Figure 4.1: VLL Example

To a customer it looks the same as if a single ATM VCC or Frame Relay circuit were used to interconnect the CPE devices, and the customer could be unaware that part of the circuit was in fact implemented over an IP backbone. This may be useful, for example, if a provider wishes to provide a LAN interconnect service using ATM as the network interface, but does not have an ATM network that directly interconnects all possible customer sites.

It is not necessary that the two links used to connect the CPE devices to the ISP nodes be of the same media type, but in this case the ISP nodes cannot treat the traffic in an opaque manner, as described above. Instead the ISP nodes must perform the functions of an interworking device between the two media types (e.g., ATM and Frame Relay), and perform functions such as LLC/SNAP to NLPID conversion, mapping between ARP protocol variants and performing any media specific processing that may be expected by the CPE devices (e.g., ATM OAM cell handling or Frame Relay XID exchanges).

The IP tunneling protocol used must support multiprotocol operation and may need to support sequencing, if that characteristic is important to the customer traffic. If the tunnels are established using a signalling protocol, they may be set up in a data driven manner, when a frame is received from a customer link and no tunnel exists, or the tunnels may be established at provisioning time and kept up permanently.

Note that the use of the term 'VLL' in this document is different to that used in the definition of the Diffserv Expedited Forwarding Per Hop Behaviour (EF-PHB) [30]. In that document a VLL is used to mean a low latency, low jitter, assured bandwidth path, which can be provided using the described PHB. Thus the focus there is primarily on link characteristics that are temporal in nature. In this document the term VLL does not imply the use of any specific QoS mechanism, Diffserv or otherwise. Instead the focus is primarily on link characteristics that are more topological in nature, (e.g., such as constructing a link which includes an IP tunnel as one segment of the link). For a truly complete emulation of a link layer both the temporal and topological aspects need to be taken into account.

5.0 VPN Types: Virtual Private Routed Networks

5.1 VPRN Characteristics

A Virtual Private Routed Network (VPRN) is defined to be the emulation of a multi-site wide area routed network using IP facilities. This section looks at how a network-based VPRN service can be provided. CPE-based VPRNs are also possible, but are not specifically discussed here. With network-based VPRNs many of the issues that need to be addressed are concerned with configuration and operational issues, which must take into account the split in administrative responsibility between the service provider and the service user.

The distinguishing characteristic of a VPRN, in comparison to other types of VPNs, is that packet forwarding is carried out at the network layer. A VPRN consists of a mesh of IP tunnels between ISP routers, together with the routing capabilities needed to forward traffic received at each VPRN node to the appropriate destination site. Attached to the ISP routers are CPE routers connected via one or more links, termed 'stub' links. There is a VPRN specific forwarding table at each ISP router to which members of the VPRN are connected. Traffic is forwarded between ISP routers, and between ISP routers and customer sites, using these forwarding tables, which contain network layer reachability information (in contrast to a Virtual Private LAN Segment type of VPN (VPLS) where the forwarding tables contain MAC layer reachability information - see section 7.0).

An example VPRN is illustrated in the following diagram, which shows 3 ISP edge routers connected via a full mesh of IP tunnels, used to interconnect 4 CPE routers. One of the CPE routers is multihomed to the ISP network. In the multihomed case, all stub links may be active, or, as shown, there may be one primary and one or more backup links to be used in case of failure of the primary. The term 'backdoor' link is used to refer to a link between two customer sites

that does not traverse the ISP network.

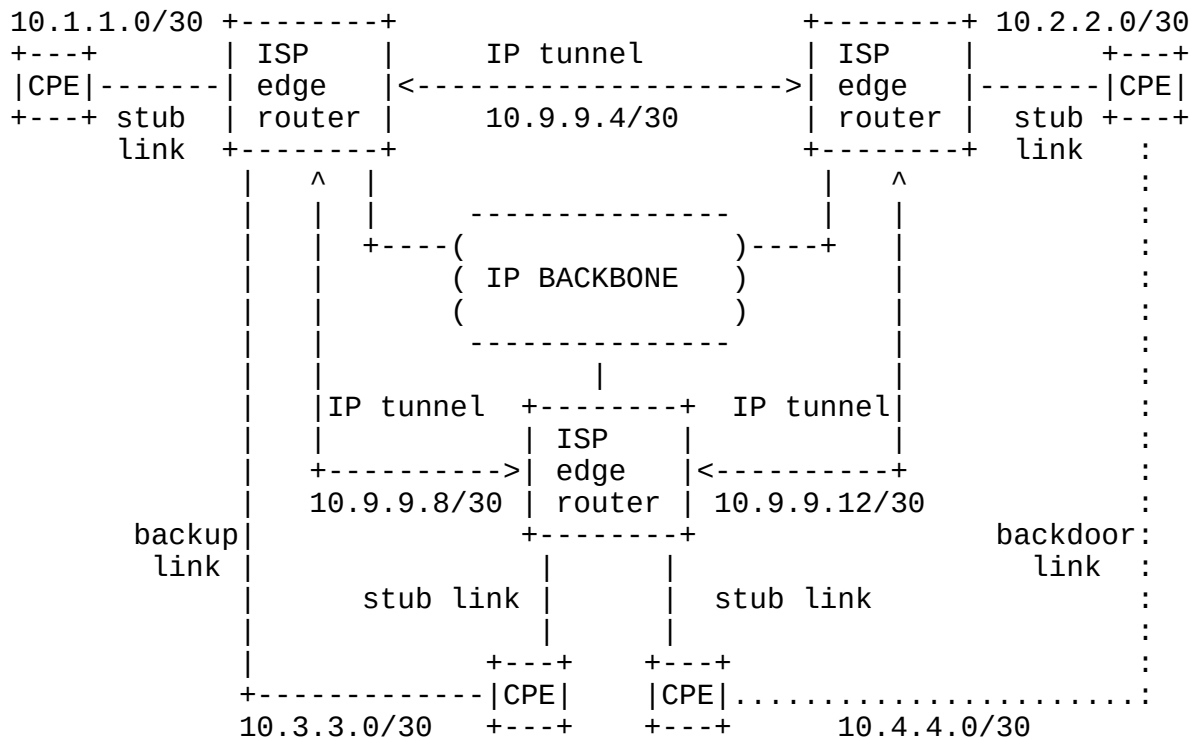


Figure 5.1: VPRN Example

The principal benefit of a VPRN is that the complexity and the configuration of the CPE routers is minimized. To a CPE router, the ISP edge router appears as a neighbor router in the customer's network, to which it sends all traffic, using a default route. The tunnel mesh that is set up to transfer traffic extends between the ISP edge routers, not the CPE routers. In effect the burden of tunnel establishment and maintenance and routing configuration is outsourced to the ISP. In addition other services needed for the operation of a VPN such as the provision of a firewall and QoS processing can be handled by a small number of ISP edge routers, rather than a large number of potentially heterogeneous CPE devices. The introduction and management of new services can also be more easily handled, as this can be achieved without the need to upgrade any CPE equipment. This latter benefit is particularly important when there may be large numbers of residential subscribers using VPN services to access private corporate networks. In this respect the model is somewhat akin to that used for telephony services, whereby new services (e.g., call waiting) can be introduced with no change in subscriber equipment.

The VPRN type of VPN is in contrast to one where the tunnel mesh extends to the CPE routers, and where the ISP network provides layer 2 connectivity alone. The latter case can be implemented either as a set of VLLs between CPE routers (see section 4.0), in which case the ISP network provides a set of layer 2 point-to-point links, or as a VPLS (see section 7.0), in which case the ISP network is used to emulate a multiaccess LAN segment. With these scenarios a customer may have more flexibility (e.g., any IGP or any protocol can be run across all customer sites) but this usually comes at the expense of a more complex configuration for the customer. Thus, depending on customer requirements, a VPRN or a VPLS may be the more appropriate solution.

Because a VPRN carries out forwarding at the network layer, a single VPRN only directly supports a single network layer protocol. For multiprotocol support, a separate VPRN for each network layer protocol could be used, or one protocol could be tunneled over another (e.g., non-IP protocols tunneled over an IP VPRN) or alternatively the ISP network could be used to provide layer 2 connectivity only, such as with a VPLS as mentioned above.

The issues to be addressed for VPRNs include initial configuration, determination by an ISP edge router of the set of links that are in each VPRN, the set of other routers that have members in the VPRN, and the set of IP address prefixes reachable via each stub link, determination by a CPE router of the set of IP address prefixes to be forwarded to an ISP edge router, the mechanism used to disseminate stub reachability information to the correct set of ISP routers, and the establishment and use of the tunnels used to carry the data traffic. Note also that, although discussed first for VPRNs, many of these issues also apply to the VPLS scenario described later, with the network layer addresses being replaced by link layer addresses.

Note that VPRN operation is decoupled from the mechanisms used by the customer sites to access the Internet. A typical scenario would be for the ISP edge router to be used to provide both VPRN and Internet connectivity to a customer site. In this case the CPE router just has a default route pointing to the ISP edge router, with the latter being responsible for steering private traffic to the VPRN and other traffic to the Internet, and providing firewall functionality between the two domains. Alternatively a customer site could have Internet connectivity via an ISP router not involved in the VPRN, or even via a different ISP. In this case the CPE device is responsible for splitting the traffic into the two domains and providing firewall functionality.

5.1.1 Topology

The topology of a VPRN may consist of a full mesh of tunnels between each VPRN node, or may be an arbitrary topology, such as a set of remote offices connected to the nearest regional site, with these regional sites connected together via a full or partial mesh. With VPRNs using IP tunnels there is much less cost assumed with full meshing than in cases where physical resources (e.g., a leased line) must be allocated for each connected pair of sites, or where the tunneling method requires resources to be allocated in the devices used to interconnect the edge routers (e.g., Frame Relay DLCIs). A full mesh topology yields optimal routing, since it precludes the need for traffic between two sites to traverse a third. Another attraction of a full mesh is that there is no need to configure topology information for the VPRN. Instead, given the member routers of a VPRN, the topology is implicit. If the number of ISP edge routers in a VPRN is very large, however, a full mesh topology may not be appropriate, due to the scaling issues involved, for example, the growth in the number of tunnels needed between sites, (which for n sites is $n(n-1)/2$), or the number of routing peers per router. Network policy may also lead to non full mesh topologies, for example an administrator may wish to set up the topology so that traffic between two remote sites passes through a central site, rather than go directly between the remote sites. It is also necessary to deal with the scenario where there is only partial connectivity across the IP backbone under certain error conditions (e.g. A can reach B, and B can reach C, but A cannot reach C directly), which can occur if policy routing is being used.

For a network-based VPRN, it is assumed that each customer site CPE router connects to an ISP edge router through one or more point-to-point stub links (e.g. leased lines, ATM or Frame Relay connections). The ISP routers are responsible for learning and disseminating reachability information amongst themselves. The CPE routers must learn the set of destinations reachable via each stub link, though this may be as simple as a default route.

The stub links may either be dedicated links, set up via provisioning, or may be dynamic links set up on demand, for example using PPP, voluntary tunneling (see section 6.3), or ATM signalling. With dynamic links it is necessary to authenticate the subscriber, and determine the authorized resources that the subscriber can access (e.g. which VPRNs the subscriber may join). Other than the way the subscriber is initially bound to the VPRN, (and this process may involve extra considerations such as dynamic IP address assignment), the subsequent VPRN mechanisms and services can be used for both types of subscribers in the same way.

5.1.2 Addressing

The addressing used within a VPRN may have no relation to the addressing used on the IP backbone over which the VPRN is instantiated. In particular non-unique private IP addressing may be used [4]. Multiple VPRNs may be instantiated over the same set of physical devices, and they may use the same or overlapping address spaces.

5.1.3 Forwarding

For a VPRN the tunnel mesh forms an overlay network operating over an IP backbone. Within each of the ISP edge routers there must be VPN specific forwarding state to forward packets received from stub links ('ingress traffic') to the appropriate next hop router, and to forward packets received from the core ('egress traffic') to the appropriate stub link. For cases where an ISP edge router supports multiple stub links belonging to the same VPRN, the tunnels can, as a local matter, either terminate on the edge router, or on a stub link. In the former case a VPN specific forwarding table is needed for egress traffic, in the latter case it is not. A VPN specific forwarding table is generally needed in the ingress direction, in order to direct traffic received on a stub link onto the correct IP tunnel towards the core.

Also since a VPRN operates at the internetwork layer, the IP packets sent over a tunnel will have their Time to Live (TTL) field decremented in the normal manner, preventing packets circulating indefinitely in the event of a routing loop within the VPRN.

5.1.4 Multiple concurrent VPRN connectivity

Note also that a single customer site may belong concurrently to multiple VPRNs and may want to transmit traffic both onto one or more VPRNs and to the default Internet, over the same stub link. There are a number of possible approaches to this problem, but these are outside the scope of this document.

5.2 VPRN Related Work

VPRN requirements and mechanisms have been discussed previously in a number of different documents. One of the first was [10], which showed how the same VPN functionality can be implemented over both MPLS and non-MPLS networks. Some others are briefly discussed below.

There are two main variants as regards the mechanisms used to provide VPRN membership and reachability functionality, - overlay and piggybacking. These are discussed in greater detail in sections

5.3.2, 5.3.3 and 5.3.4 below. An example of the overlay model is described in [14], which discusses the provision of VPRN functionality by means of a separate per-VPN routing protocol instance and route and forwarding table instantiation, otherwise known as virtual routing. Each VPN routing instance is isolated from any other VPN routing instance, and from the routing used across the backbone. As a result any routing protocol (e.g. OSPF, RIP2, IS-IS) can be run with any VPRN, independently of the routing protocols used in other VPRNs, or in the backbone itself. The VPN model described in [12] is also an overlay VPRN model using virtual routing. That document is specifically geared towards the provision of VPRN functionality over MPLS backbones, and it describes how VPRN membership dissemination can be automated over an MPLS backbone, by performing VPN neighbor discovery over the base MPLS tunnel mesh. [31] extends the virtual routing model to include VPN areas, and VPN border routers which route between VPN areas. VPN areas may be defined for administrative or technical reasons, such as different underlying network infrastructures (e.g. ATM, MPLS, IP).

In contrast [15] describes the provision of VPN functionality using a piggybacking approach for membership and reachability dissemination, with this information being piggybacked in Border Gateway Protocol 4 (BGP) [32] packets. VPNs are constructed using BGP policies, which are used to control which sites can communicate with each other. [13] also uses BGP for piggybacking membership information, and piggybacks reachability information on the protocol used to establish MPLS LSPs (CR-LDP or extended RSVP). Unlike the other proposals, however, this proposal requires the participation on the CPE router to implement the VPN functionality.

5.3 VPRN Generic Requirements

There are a number of common requirements which any network-based VPRN solution must address, and there are a number of different mechanisms that can be used to meet these requirements. These generic issues are

- 1) The use of a globally unique VPN identifier in order to be able to refer to a particular VPN.
- 2) VPRN membership determination. An edge router must learn of the local stub links that are in each VPRN, and must learn of the set of other routers that have members in that VPRN.
- 3) Stub link reachability information. An edge router must learn the set of addresses and address prefixes reachable via each stub link.

- 4) Intra-VPN reachability information. Once an edge router has determined the set of address prefixes associated with each of its stub links, then this information must be disseminated to each other edge router in the VPN.
- 5) Tunneling mechanism. An edge router must construct the necessary tunnels to other routers that have members in the VPN, and must perform the encapsulation and decapsulation necessary to send and receive packets over the tunnels.

5.3.1 VPN Identifier

The IETF [16] and the ATM Forum [17] have standardized on a single format for a globally unique identifier used to identify a VPN - a VPN-ID. Only the format of the VPN-ID has been defined, not its semantics or usage. The aim is to allow its use for a wide variety of purposes, and to allow the same identifier to be used with different technologies and mechanisms. For example a VPN-ID can be included in a MIB to identify a VPN for management purposes. A VPN-ID can be used in a control plane protocol, for example to bind a tunnel to a VPN at tunnel establishment time. All packets that traverse the tunnel are then implicitly associated with the identified VPN. A VPN-ID can be used in a data plane encapsulation, to allow for an explicit per-packet identification of the VPN associated with the packet. If a VPN is implemented using different technologies (e.g., IP and ATM) in a network, the same identifier can be used to identify the VPN across the different technologies. Also if a VPN spans multiple administrative domains the same identifier can be used everywhere.

Most of the VPN schemes developed (e.g. [11], [12], [13], [14]) require the use of a VPN-ID that is carried in control and/or data packets, which is used to associate the packet with a particular VPN. Although the use of a VPN-ID in this manner is very common, it is not universal. [15] describes a scheme where there is no protocol field used to identify a VPN in this manner. In this scheme the VPNs as understood by a user, are administrative constructs, built using BGP policies. There are a number of attributes associated with VPN routes, such as a route distinguisher, and origin and target "VPN", that are used by the underlying protocol mechanisms for disambiguation and scoping, and these are also used by the BGP policy mechanism in the construction of VPNs, but there is nothing corresponding with the VPN-ID as used in the other documents.

Note also that [33] defines a multiprotocol encapsulation for use over ATM AAL5 that uses the standard VPN-ID format.

5.3.2 VPN Membership Information Configuration and Dissemination

In order to establish a VPRN, or to insert new customer sites into an established VPRN, an ISP edge router must determine which stub links are associated with which VPRN. For static links (e.g. an ATM VCC) this information must be configured into the edge router, since the edge router cannot infer such bindings by itself. An SNMP MIB allowing for bindings between local stub links and VPN identities is one solution.

For subscribers that attach to the network dynamically (e.g. using PPP or voluntary tunneling) it is possible to make the association between stub link and VPRN as part of the end user authentication processing that must occur with such dynamic links. For example the VPRN to which a user is to be bound may be derived from the domain name the user used as part of PPP authentication. If the user is successfully authenticated (e.g. using a Radius server), then the newly created dynamic link can be bound to the correct VPRN. Note that static configuration information is still needed, for example to maintain the list of authorized subscribers for each VPRN, but the location of this static information could be an external authentication server rather than on an ISP edge router. Whether the link was statically or dynamically created, a VPN-ID can be associated with that link to signify to which VPRN it is bound.

After learning which stub links are bound to which VPRN, each edge router must learn either the identity of, or, at least, the route to, each other edge router supporting other stub links in that particular VPRN. Implicit in the latter is the notion that there exists some mechanism by which the configured edge routers can then use this edge router and/or stub link identity information to subsequently set up the appropriate tunnels between them. The problem of VPRN member dissemination between participating edge routers, can be solved in a variety of ways, discussed below.

5.3.2.1 Directory Lookup

The members of a particular VPRN, that is, the identity of the edge routers supporting stub links in the VPRN, and the set of static stub links bound to the VPRN per edge router, could be configured into a directory, which edge routers could query, using some defined mechanism (e.g. Lightweight Directory Access Protocol (LDAP) [34]), upon startup.

Using a directory allows either a full mesh topology or an arbitrary topology to be configured. For a full mesh, the full list of member routers in a VPRN is distributed everywhere. For an arbitrary topology, different routers may receive different member lists.

Using a directory allows for authorization checking prior to disseminating VPRN membership information, which may be desirable where VPRNs span multiple administrative domains. In such a case, directory to directory protocol mechanisms could also be used to propagate authorized VPRN membership information between the directory systems of the multiple administrative domains.

There also needs to be some form of database synchronization mechanism (e.g. triggered or regular polling of the directory by edge routers, or active pushing of update information to the edge routers by the directory) in order for all edge routers to learn the identity of newly configured sites inserted into an active VPRN, and also to learn of sites removed from a VPRN.

5.3.2.2 Explicit Management Configuration

A VPRN MIB could be defined which would allow a central management system to configure each edge router with the identities of each other participating edge router and the identity of each of the static stub links bound to the VPRN. Like the use of a directory, this mechanism allows both full mesh and arbitrary topologies to be configured. Another mechanism using a centralized management system is to use a policy server and use the Common Open Policy Service (COPS) protocol [35] to distribute VPRN membership and policy information, such as the tunnel attributes to use when establishing a tunnel, as described in [36].

Note that this mechanism allows the management station to impose strict authorization control; on the other hand, it may be more difficult to configure edge routers outside the scope of the management system. The management configuration model can also be considered a subset of the directory method, in that the management directories could use MIBs to push VPRN membership information to the participating edge routers, either subsequent to, or as part of, the local stub link configuration process.

5.3.2.3 Piggybacking in Routing Protocols

VPRN membership information could be piggybacked into the routing protocols run by each edge router across the IP backbone, since this is an efficient means of automatically propagating information throughout the network to other participating edge routers. Specifically, each route advertisement by each edge router could include, at a minimum, the set of VPN identifiers associated with each edge router, and adequate information to allow other edge routers to determine the identity of, and/or, the route to, the particular edge router. Other edge routers would examine received route advertisements to determine if any contained information was

relevant to a supported (i.e., configured) VPRN; this determination could be done by looking for a VPN identifier matching a locally configured VPN. The nature of the piggybacked information, and related issues, such as scoping, and the means by which the nodes advertising particular VPN memberships will be identified, will generally be a function both of the routing protocol and of the nature of the underlying transport.

Using this method all the routers in the network will have the same view of the VPRN membership information, and so a full mesh topology is easily supported. Supporting an arbitrary topology is more difficult, however, since some form of pruning would seem to be needed.

The advantage of the piggybacking scheme is that it allows for efficient information dissemination, but it does require that all nodes in the path, and not just the participating edge routers, be able to accept such modified route advertisements. A disadvantage is that significant administrative complexity may be required to configure scoping mechanisms so as to both permit and constrain the dissemination of the piggybacked advertisements, and in itself this may be quite a configuration burden, particularly if the VPRN spans multiple routing domains (e.g. different autonomous systems / ISPs).

Furthermore, unless some security mechanism is used for routing updates so as to permit only all relevant edge routers to read the piggybacked advertisements, this scheme generally implies a trust model where all routers in the path must perforce be authorized to know this information. Depending upon the nature of the routing protocol, piggybacking may also require intermediate routers, particularly autonomous system (AS) border routers, to cache such advertisements and potentially also re-distribute them between multiple routing protocols.

Each of the schemes described above have merit in particular situations. Note that, in practice, there will almost always be some centralized directory or management system which will maintain VPRN membership information, such as the set of edge routers that are allowed to support a certain VPRN, the bindings of static stub links to VPRNs, or authentication and authorization information for users that access the network via dynamics links. This information needs to be configured and stored in some form of database, so that the additional steps needed to facilitate the configuration of such information into edge routers, and/or, facilitate edge router access to such information, may not be excessively onerous.

5.3.3 Stub Link Reachability Information

There are two aspects to stub site reachability - the means by which VPRN edge routers determine the set of VPRN addresses and address prefixes reachable at each stub site, and the means by which the CPE routers learn the destinations reachable via each stub link. A number of common scenarios are outlined below. In each case the information needed by the ISP edge router is the same - the set of VPRN addresses reachable at the customer site, but the information needed by the CPE router differs.

5.3.3.1 Stub Link Connectivity Scenarios

5.3.3.1.1 Dual VPRN and Internet Connectivity

The CPE router is connected via one link to an ISP edge router, which provides both VPRN and Internet connectivity.

This is the simplest case for the CPE router, as it just needs a default route pointing to the ISP edge router.

5.3.3.1.2 VPRN Connectivity Only

The CPE router is connected via one link to an ISP edge router, which provides VPRN, but not Internet, connectivity.

The CPE router must know the set of non-local VPRN destinations reachable via that link. This may be a single prefix, or may be a number of disjoint prefixes. The CPE router may be either statically configured with this information, or may learn it dynamically by running an instance of an Interior Gateway Protocol (IGP). For simplicity it is assumed that the IGP used for this purpose is RIP, though it could be any IGP. The ISP edge router will inject into this instance of RIP the VPRN routes which it learns by means of one of the intra-VPRN reachability mechanisms described in section 5.3.4. Note that the instance of RIP run to the CPE, and any instance of a routing protocol used to learn intra-VPRN reachability (even if also RIP) are separate, with the ISP edge router redistributing the routes from one instance to another.

5.3.3.1.3 Multihomed Connectivity

The CPE router is multihomed to the ISP network, which provides VPRN connectivity.

In this case all the ISP edge routers could advertise the same VPRN routes to the CPE router, which then sees all VPRN prefixes equally reachable via all links. More specific route redistribution is also possible, whereby each ISP edge router advertises a different set of prefixes to the CPE router.

5.3.3.1.4 Backdoor Links

The CPE router is connected to the ISP network, which provides VPRN connectivity, but also has a backdoor link to another customer site

In this case the ISP edge router will advertise VPRN routes as in case 2 to the CPE device. However now the same destination is reachable via both the ISP edge router and via the backdoor link. If the CPE routers connected to the backdoor link are running the customer's IGP, then the backdoor link may always be the favored link as it will appear as an 'internal' path, whereas the destination as injected via the ISP edge router will appear as an 'external' path (to the customer's IGP). To avoid this problem, assuming that the customer wants the traffic to traverse the ISP network, then a separate instance of RIP should be run between the CPE routers at both ends of the backdoor link, in the same manner as an instance of RIP is run on a stub or backup link between a CPE router and an ISP edge router. This will then also make the backdoor link appear as an external path, and by adjusting the link costs appropriately, the ISP path can always be favored, unless it goes down, when the backdoor link is then used.

The description of the above scenarios covers what reachability information is needed by the ISP edge routers and the CPE routers, and discusses some of the mechanisms used to convey this information. The sections below look at these mechanisms in more detail.

5.3.3.1 Routing Protocol Instance

A routing protocol can be run between the CPE edge router and the ISP edge router to exchange reachability information. This allows an ISP edge router to learn the VPRN prefixes reachable at a customer site, and also allows a CPE router to learn the destinations reachable via the provider network.

The extent of the routing domain for this protocol instance is generally just the ISP edge router and the CPE router although if the customer site is also running the same protocol as its IGP, then the domain may extend into customer site. If the customer site is running a different routing protocol then the CPE router redistributes the routes between the instance running to the ISP edge router, and the instance running into the customer site.

Given the typically restricted scope of this routing instance, a simple protocol will generally suffice. RIP is likely to be the most common protocol used, though any routing protocol, such as OSPF, or BGP run in internal mode (IBGP), could also be used.

Note that the instance of the stub link routing protocol is different from any instance of a routing protocol used for intra-VPN reachability. For example, if the ISP edge router uses routing protocol piggybacking to disseminate VPN membership and reachability information across the core, then it may redistribute suitably labeled routes from the CPE routing instance to the core routing instance. The routing protocols used for each instance are decoupled, and any suitable protocol can be used in each case. There is no requirement that the same protocol, or even the same stub link reachability information gathering mechanism, be run between each CPE router and associated ISP edge router in a particular VPN, since this is a purely local matter.

This decoupling allows ISPs to deploy a common (across all VPNs) intra-VPN reachability mechanism, and a common stub link reachability mechanism, with these mechanisms isolated both from each other, and from the particular IGP used in a customer network. In the first case, due to the IGP-IGP boundary implemented on the ISP edge router, the ISP can insulate the intra-VPN reachability mechanism from misbehaving stub link protocol instances. In the second case the ISP is not required to be aware of the particular IGP running in a customer site. Other scenarios are possible, where the ISP edge routers are running a routing protocol in the same instance as the customer's IGP, but are unlikely to be practical, since it defeats the purpose of a VPN simplifying CPE router configuration. In cases where a customer wishes to run an IGP across multiple sites, a VPLS solution is more suitable.

Note that if a particular customer site concurrently belongs to multiple VPNs (or wishes to concurrently communicate with both a VPN and the Internet), then the ISP edge router must have some means of unambiguously mapping stub link address prefixes to particular VPNs. A simple way is to have multiple stub links, one per VPN. It is also possible to run multiple VPNs over one stub link. This could be done either by ensuring (and appropriately configuring the

ISP edge router to know) that particular disjoint address prefixes are mapped into separate VPRNs, or by tagging the routing advertisements from the CPE router with the appropriate VPN identifier. For example if MPLS was being used to convey stub link reachability information, different MPLS labels would be used to differentiate the disjoint prefixes assigned to particular VPRNs. In any case, some administrative procedure would be required for this coordination.

5.3.3.2 Configuration

The reachability information across each stub link could be manually configured, which may be appropriate if the set of addresses or prefixes is small and static.

5.3.3.3 ISP Administered Addresses

The set of addresses used by each stub site could be administered and allocated via the VPRN edge router, which may be appropriate for small customer sites, typically containing either a single host, or a single subnet. Address allocation can be carried out using protocols such as PPP or DHCP [37], with, for example, the edge router acting as a Radius client and retrieving the customer's IP address to use from a Radius server, or acting as a DHCP relay and examining the DHCP reply message as it is relayed to the customer site. In this manner the edge router can build up a table of stub link reachability information. Although these address assignment mechanisms are typically used to assign an address to a single host, some vendors have added extensions whereby an address prefix can be assigned, with, in some cases, the CPE device acting as a "mini-DHCP" server and assigning addresses for the hosts in the customer site.

Note that with these schemes it is the responsibility of the address allocation server to ensure that each site in the VPN received a disjoint address space. Note also that an ISP would typically only use this mechanism for small stub sites, which are unlikely to have backdoor links.

5.3.3.4 MPLS Label Distribution Protocol

In cases where the CPE router runs MPLS, LDP can be used to convey the set of prefixes at a stub site to a VPRN edge router. Using the downstream unsolicited mode of label distribution the CPE router can distribute a label for each route in the stub site. Note however that the processing carried out by the edge router in this case is more than just the normal LDP processing, since it is learning new routes via LDP, rather than the usual case of learning labels for existing routes that it has learned via standard routing mechanisms.

5.3.4 Intra-VPN Reachability Information

Once an edge router has determined the set of prefixes associated with each of its stub links, then this information must be disseminated to each other edge router in the VPRN. Note also that there is an implicit requirement that the set of reachable addresses within the VPRN be locally unique that is, each VPRN stub link (not performing load sharing) maintain an address space disjoint from any other, so as to permit unambiguous routing. In practical terms, it is also generally desirable, though not required, that this address space be well partitioned i.e., specific, disjoint address prefixes per edge router, so as to preclude the need to maintain and disseminate large numbers of host routes.

The problem of intra-VPN reachability information dissemination can be solved in a number of ways, some of which include the following:

5.3.4.1 Directory Lookup

Along with VPRN membership information, a central directory could maintain a listing of the address prefixes associated with each customer site. Such information could be obtained by the server through protocol interactions with each edge router. Note that the same directory synchronization issues discussed above in section 5.3.2 also apply in this case.

5.3.4.2 Explicit Configuration

The address spaces associated with each edge router could be explicitly configured into each other router. This is clearly a non-scalable solution, particularly when arbitrary topologies are used, and also raises the question of how the management system learns such information in the first place.

5.3.4.3 Local Intra-VPN Routing Instantiations

In this approach, each edge router runs an instance of a routing protocol (a 'virtual router') per VPRN, running across the VPRN tunnels to each peer edge router, to disseminate intra-VPN reachability information. Both full-mesh and arbitrary VPRN topologies can be easily supported, since the routing protocol itself can run over any topology. The intra-VPN routing advertisements could be distinguished from normal tunnel data packets either by being addressed directly to the peer edge router, or by a tunnel specific mechanism.

Note that this intra-VPN routing protocol need have no relationship either with the IGP of any customer site or with the routing protocols operated by the ISPs in the IP backbone. Depending on the size and scale of the VPNs to be supported either a simple protocol like RIP or a more sophisticated protocol like OSPF could be used. Because the intra-VPN routing protocol operates as an overlay over the IP backbone it is wholly transparent to any intermediate routers, and to any edge routers not within the VPN. This also implies that such routing information can remain opaque to such routers, which may be a necessary security requirements in some cases. Also note that if the routing protocol runs directly over the same tunnels as the data traffic, then it will inherit the same level of security as that afforded the data traffic, for example strong encryption and authentication.

If the tunnels over which an intra-VPN routing protocol runs are dedicated to a specific VPN (e.g. a different multiplexing field is used for each VPN) then no changes are needed to the routing protocol itself. On the other hand if shared tunnels are used, then it is necessary to extend the routing protocol to allow a VPN-ID field to be included in routing update packets, to allow sets of prefixes to be associated with a particular VPN.

5.3.4.4 Link Reachability Protocol

By link reachability protocol is meant a protocol that allows two nodes, connected via a point-to-point link, to exchange reachability information. Given a full mesh topology, each edge router could run a link reachability protocol, for instance some variation of MPLS CR-LDP, across the tunnel to each peer edge router in the VPN, carrying the VPN-ID and the reachability information of each VPN running across the tunnel between the two edge routers. If VPN membership information has already been distributed to an edge router, then the neighbor discovery aspects of a traditional routing protocol are not needed, as the set of neighbors is already known. TCP connections can be used to interconnect the neighbors, to provide reliability. This approach may reduce the processing burden of running routing protocol instances per VPN, and may be of particular benefit where a shared tunnel mechanism is used to connect a set of edge routers supporting multiple VPNs.

Another approach to developing a link reachability protocol would be to base it on IBGP. The problem that needs to be solved by a link reachability protocol is very similar to that solved by IBGP - conveying address prefixes reliably between edge routers.

Using a link reachability protocol it is straightforward to support a full mesh topology - each edge router conveys its own local reachability information to all other routers, but does not redistribute information received from any other router. However once an arbitrary topology needs to be supported, the link reachability protocol needs to develop into a full routing protocol, due to the need to implement mechanisms to avoid loops, and there would seem little benefit in reinventing another routing protocol to deal with this. Some reasons why partially connected meshes may be needed even in a tunneled environment are discussed in section 5.1.1.

5.3.4.5 Piggybacking in IP Backbone Routing Protocols

As with VPRN membership, the set of address prefixes associated with each stub interface could also be piggybacked into the routing advertisements from each edge router and propagated through the network. Other edge routers extract this information from received route advertisements in the same way as they obtain the VPRN membership information (which, in this case, is implicit in the identification of the source of each route advertisement). Note that this scheme may require, depending upon the nature of the routing protocols involved, that intermediate routers, e.g. border routers, cache intra-VPRN routing information in order to propagate it further. This also has implications for the trust model, and for the level of security possible for intra-VPRN routing information.

Note that in any of the cases discussed above, an edge router has the option of disseminating its stub link prefixes in a manner so as to permit tunneling from remote edge routers directly to the egress stub links. Alternatively, it could disseminate the information so as to associate all such prefixes with the edge router, rather than with specific stub links. In this case, the edge router would need to implement a VPN specific forwarding mechanism for egress traffic, to determine the correct egress stub link. The advantage of this is that it may significantly reduce the number of distinct tunnels or tunnel label information which need to be constructed and maintained. Note that this choice is purely a local manner and is not visible to remote edge routers.

5.3.5 Tunneling Mechanisms

Once VPRN membership information has been disseminated, the tunnels comprising the VPRN core can be constructed.

One approach to setting up the tunnel mesh is to use point-to-point IP tunnels, and the requirements and issues for such tunnels have been discussed in section 3.0. For example while tunnel establishment can be done through manual configuration, this is

clearly not likely to be a scalable solution, given the $O(n^2)$ problem of meshed links. As such, tunnel set up should use some form of signalling protocol to allow two nodes to construct a tunnel to each other knowing only each other's identity.

Another approach is to use the multipoint to point 'tunnels' provided by MPLS. As noted in [38], MPLS can be considered to be a form of IP tunneling, since the labels of MPLS packets allow for routing decisions to be decoupled from the addressing information of the packets themselves. MPLS label distribution mechanisms can be used to associate specific sets of MPLS labels with particular VPRN address prefixes supported on particular egress points (i.e., stub links of edge routers) and hence allow other edge routers to explicitly label and route traffic to particular VPRN stub links.

One attraction of MPLS as a tunneling mechanism is that it may require less processing within each edge router than alternative tunneling mechanisms. This is a function of the fact that data security within a MPLS network is implicit in the explicit label binding, much as with a connection oriented network, such as Frame Relay. This may hence lessen customer concerns about data security and hence require less processor intensive security mechanisms (e.g., IPsec). However there are other potential security concerns with MPLS. There is no direct support for security features such as authentication, confidentiality, and non-repudiation and the trust model for MPLS means that intermediate routers, (which may belong to different administrative domains), through which membership and prefix reachability information is conveyed, must be trusted, not just the edge routers themselves.

5.4 Multihomed Stub Routers

The discussion thus far has implicitly assumed that stub routers are connected to one and only one VPRN edge router. In general, this restriction should be capable of being relaxed without any change to VPRN operation, given general market interest in multihoming for reliability and other reasons. In particular, in cases where the stub router supports multiple redundant links, with only one operational at any given time, with the links connected either to the same VPRN edge router, or to two or more different VPRN edge routers, then the stub link reachability mechanisms will both discover the loss of an active link, and the activation of a backup link. In the former situation, the previously connected VPRN edge router will cease advertising reachability to the stub node, while the VPRN edge router with the now active link will begin advertising reachability, hence restoring connectivity.

An alternative scenario is where the stub node supports multiple active links, using some form of load sharing algorithm. In such a case, multiple VPRN edge routers may have active paths to the stub node, and may so advertise across the VPRN. This scenario should not cause any problem with reachability across the VPRN providing that the intra-VPRN reachability mechanism can accommodate multiple paths to the same prefix, and has the appropriate mechanisms to preclude looping - for instance, distance vector metrics associated with each advertised prefix.

5.5 Multicast Support

Multicast and broadcast traffic can be supported across VPRNs either by edge replication or by native multicast support in the backbone. These two cases are discussed below.

5.5.1 Edge Replication

This is where each VPRN edge router replicates multicast traffic for transmission across each link in the VPRN. Note that this is the same operation that would be performed by CPE routers terminating actual physical links or dedicated connections. As with CPE routers, multicast routing protocols could also be run on each VPRN edge router to determine the distribution tree for multicast traffic and hence reduce unnecessary flood traffic. This could be done by running instances of standard multicast routing protocols, e.g. Protocol Independent Multicast (PIM) [39] or Distance Vector Multicast Routing Protocol (DVMRP) [40], on and between each VPRN edge router, through the VPRN tunnels, in the same way that unicast routing protocols might be run at each VPRN edge router to determine intra-VPN unicast reachability, as discussed in section 5.3.4. Alternatively, if a link reachability protocol was run across the VPRN tunnels for intra-VPRN reachability, then this could also be augmented to allow VPRN edge routers to indicate both the particular multicast groups requested for reception at each edge node, and also the multicast sources at each edge site.

In either case, there would need to be some mechanism to allow for the VPRN edge routers to determine which particular multicast groups were requested at each site and which sources were present at each site. How this could be done would, in general, be a function of the capabilities of the CPE stub routers at each site. If these run multicast routing protocols, then they can interact directly with the equivalent protocols at each VPRN edge router. If the CPE device does not run a multicast routing protocol, then in the absence of Internet Group Management Protocol (IGMP) proxying [41] the customer site would be limited to a single subnet connected to the VPRN edge router via a bridging device, as the scope of an IGMP message is

limited to a single subnet. However using IGMP-proxying the CPE router can engage in multicast forwarding without running a multicast routing protocol, in constrained topologies. On its interfaces into the customer site the CPE router performs the router functions of IGMP, and on its interface to the VPRN edge router it performs the host functions of IGMP.

5.5.2 Native Multicast Support

This is where VPRN edge routers map intra-VPRN multicast traffic onto a native IP multicast distribution mechanism across the backbone. Note that intra-VPRN multicast has the same requirements for isolation from general backbone traffic as intra-VPRN unicast traffic. Currently the only IP tunneling mechanism that has native support for multicast is MPLS. On the other hand, while MPLS supports native transport of IP multicast packets, additional mechanisms would be needed to leverage these mechanisms for the support of intra-VPRN multicast.

For instance, each VPRN router could prefix multicast group addresses within each VPRN with the VPN-ID of that VPRN and then redistribute these, essentially treating this VPN-ID/intra-VPRN multicast address tuple as a normal multicast address, within the backbone multicast routing protocols, as with the case of unicast reachability, as discussed previously. The MPLS multicast label distribution mechanisms could then be used to set up the appropriate multicast LSPs to interconnect those sites within each VPRN supporting particular multicast group addresses. Note, however, that this would require each of the intermediate LSRs to not only be aware of each intra-VPRN multicast group, but also to have the capability of interpreting these modified advertisements. Alternatively, mechanisms could be defined to map intra-VPRN multicast groups into backbone multicast groups.

Other IP tunneling mechanisms do not have native multicast support. It may prove feasible to extend such tunneling mechanisms by allocating IP multicast group addresses to the VPRN as a whole and hence distributing intra-VPRN multicast traffic encapsulated within backbone multicast packets. Edge VPRN routers could filter out unwanted multicast groups. Alternatively, mechanisms could also be defined to allow for allocation of backbone multicast group addresses for particular intra-VPRN multicast groups, and to then utilize these, through backbone multicast protocols, as discussed above, to limit forwarding of intra-VPRN multicast traffic only to those nodes within the group.

A particular issue with the use of native multicast support is the provision of security for such multicast traffic. Unlike the case of edge replication, which inherits the security characteristics of the underlying tunnel, native multicast mechanisms will need to use some form of secure multicast mechanism. The development of architectures and solutions for secure multicast is an active research area, for example see [42] and [43]. The Secure Multicast Group (SMuG) of the IRTF has been set up to develop prototype solutions, which would then be passed to the IETF IPsec working group for standardization.

However considerably more development is needed before scalable secure native multicast mechanisms can be generally deployed.

5.6 Recommendations

The various proposals that have been developed to support some form of VPRN functionality can be broadly classified into two groups - those that utilize the router piggybacking approach for distributing VPN membership and/or reachability information ([13],[15]) and those that use the virtual routing approach ([12],[14]). In some cases the mechanisms described rely on the characteristics of a particular infrastructure (e.g. MPLS) rather than just IP.

Within the context of the virtual routing approach it may be useful to develop a membership distribution protocol based on a directory or MIB. When combined with the protocol extensions for IP tunneling protocols outlined in section 3.2, this would then provide the basis for a complete set of protocols and mechanisms that support interoperable VPRNs that span multiple administrations over an IP backbone. Note that the other major pieces of functionality needed - the learning and distribution of customer reachability information, can be performed by instances of standard routing protocols, without the need for any protocol extensions.

Also for the constrained case of a full mesh topology, the usefulness of developing a link reachability protocol could be examined, however the limitations and scalability issues associated with this topology may not make it worthwhile to develop something specific for this case, as standard routing will just work.

Extending routing protocols to allow a VPN-ID to be carried in routing update packets could also be examined, but is not necessary if VPN specific tunnels are used.

6.0 VPN Types: Virtual Private Dial Networks

A Virtual Private Dial Network (VPDN) allows for a remote user to connect on demand through an ad hoc tunnel into another site. The user is connected to a public IP network via a dial-up PSTN or ISDN link, and user packets are tunneled across the public network to the desired site, giving the impression to the user of being 'directly' connected into that site. A key characteristic of such ad hoc connections is the need for user authentication as a prime requirement, since anyone could potentially attempt to gain access to such a site using a switched dial network.

Today many corporate networks allow access to remote users through dial connections made through the PSTN, with users setting up PPP connections across an access network to a network access server, at which point the PPP sessions are authenticated using AAA systems running such standard protocols as Radius [44]. Given the pervasive deployment of such systems, any VPDN system must in practice allow for the near transparent re-use of such existing systems.

The IETF have developed the Layer 2 Tunneling Protocol (L2TP) [8] which allows for the extension of user PPP sessions from an L2TP Access Concentrator (LAC) to a remote L2TP Network Server (LNS). The L2TP protocol itself was based on two earlier protocols, the Layer 2 Forwarding protocol (L2F) [45], and the Point-to-Point Tunneling Protocol (PPTP) [46], and this is reflected in the two quite different scenarios for which L2TP can be used - compulsory tunneling and voluntary tunneling, discussed further below in sections 6.2 and 6.3.

This document focuses on the use of L2TP over an IP network (using UDP), but L2TP may also be run directly over other protocols such as ATM or Frame Relay. Issues specifically related to running L2TP over non-IP networks, such as how to secure such tunnels, are not addressed here.

6.1 L2TP protocol characteristics

This section looks at the characteristics of the L2TP tunneling protocol using the categories outlined in section 3.0.

6.1.1 Multiplexing

L2TP has inherent support for the multiplexing of multiple calls from different users over a single link. Between the same two IP endpoints, there can be multiple L2TP tunnels, as identified by a tunnel-id, and multiple sessions within a tunnel, as identified by a session-id.

6.1.2 Signalling

This is supported via the inbuilt control connection protocol, allowing both tunnels and sessions to be established dynamically.

6.1.3 Data Security

By allowing for the transparent extension of PPP from the user, through the LAC to the LNS, L2TP allows for the use of whatever security mechanisms, with respect to both connection set up, and data transfer, may be used with normal PPP connections. However this does not provide security for the L2TP control protocol itself. In this case L2TP could be further secured by running it in combination with IPSec through IP backbones [47], [48], or related mechanisms on non-IP backbones [49].

The interaction of L2TP with AAA systems for user authentication and authorization is a function of the specific means by which L2TP is used, and the nature of the devices supporting the LAC and the LNS. These issues are discussed in depth in [50].

The means by which the host determines the correct LAC to connect to, and the means by which the LAC determines which users to further tunnel, and the LNS parameters associated with each user, are outside the scope of the operation of a VPDN, but may be addressed, for instance, by evolving Internet roaming specifications [51].

6.1.4 Multiprotocol Transport

L2TP transports PPP packets (and only PPP packets) and thus can be used to carry multiprotocol traffic since PPP itself is multiprotocol.

6.1.5 Sequencing

L2TP supports sequenced delivery of packets. This is a capability that can be negotiated at session establishment, and that can be turned on and off by an LNS during a session. The sequence number field in L2TP can also be used to provide an indication of dropped packets, which is needed by various PPP compression algorithms to operate correctly. If no compression is in use, and the LNS determines that the protocols in use (as evidenced by the PPP NCP negotiations) can deal with out of sequence packets (e.g. IP), then it may disable the use of sequencing.

6.1.6 Tunnel Maintenance

A keepalive protocol is used by L2TP in order to allow it to distinguish between a tunnel outage and prolonged periods of tunnel inactivity.

6.1.7 Large MTUs

L2TP itself has no inbuilt support for a segmentation and reassembly capability, but when run over UDP/IP IP fragmentation will take place if necessary. Note that a LAC or LNS may adjust the Maximum Receive Unit (MRU) negotiated via PPP in order to preclude fragmentation, if it has knowledge of the MTU used on the path between LAC and LNS. To this end, there is a proposal to allow the use of MTU discovery for cases where the L2TP tunnel transports IP frames [52].

6.1.8 Tunnel Overhead

L2TP as used over IP networks runs over UDP and must be used to carry PPP traffic. This results in a significant amount of overhead, both in the data plane with UDP, L2TP and PPP headers, and also in the control plane, with the L2TP and PPP control protocols. This is discussed further in section 6.3

6.1.9 Flow and Congestion Control

L2TP supports flow and congestion control mechanisms for the control protocol, but not for data traffic. See section 3.1.9 for more details.

6.1.10 QoS / Traffic Management

An L2TP header contains a 1-bit priority field, which can be set for packets that may need preferential treatment (e.g. keepalives) during local queuing and transmission. Also by transparently extending PPP, L2TP has inherent support for such PPP mechanisms as multi-link PPP [53] and its associated control protocols [54], which allow for bandwidth on demand to meet user requirements.

In addition L2TP calls can be mapped into whatever underlying traffic management mechanisms may exist in the network, and there are proposals to allow for requests through L2TP signalling for specific differentiated services behaviors [55].

6.1.11 Miscellaneous

Since L2TP is designed to transparently extend PPP, it does not attempt to supplant the normal address assignment mechanisms associated with PPP. Hence, in general terms the host initiating the PPP session will be assigned an address by the LNS using PPP procedures. This addressing may have no relation to the addressing used for communication between the LAC and LNS. The LNS will also need to support whatever forwarding mechanisms are needed to route traffic to and from the remote host.

6.2 Compulsory Tunneling

Compulsory tunneling refers to the scenario in which a network node - a dial or network access server, for instance - acting as a LAC, extends a PPP session across a backbone using L2TP to a remote LNS, as illustrated below. This operation is transparent to the user initiating the PPP session to the LAC. This allows for the decoupling of the location and/or ownership of the modem pools used to terminate dial calls, from the site to which users are provided access. Support for this scenario was the original intent of the L2F specification, upon which the L2TP specification was based.

There are a number of different deployment scenarios possible. One example, shown in the diagram below, is where a subscriber host dials into a NAS acting as a LAC, and is tunneled across an IP network (e.g. the Internet) to a gateway acting as an LNS. The gateway provides access to a corporate network, and could either be a device in the corporate network itself, or could be an ISP edge router, in the case where a customer has outsourced the maintenance of LNS functionality to an ISP. Another scenario is where an ISP uses L2TP to provide a subscriber with access to the Internet. The subscriber host dials into a NAS acting as a LAC, and is tunneled across an access network to an ISP edge router acting as an LNS. This ISP edge router then feeds the subscriber traffic into the Internet. Yet other scenarios are where an ISP uses L2TP to provide a subscriber with access to a VPRN, or with concurrent access to both a VPRN and the Internet.

A VPDN, whether using compulsory or voluntary tunneling, can be viewed as just another type of access method for subscriber traffic, and as such can be used to provide connectivity to different types of networks, e.g. a corporate network, the Internet, or a VPRN. The last scenario is also an example of how a VPN service as provided to a customer may be implemented using a combination of different types of VPN.

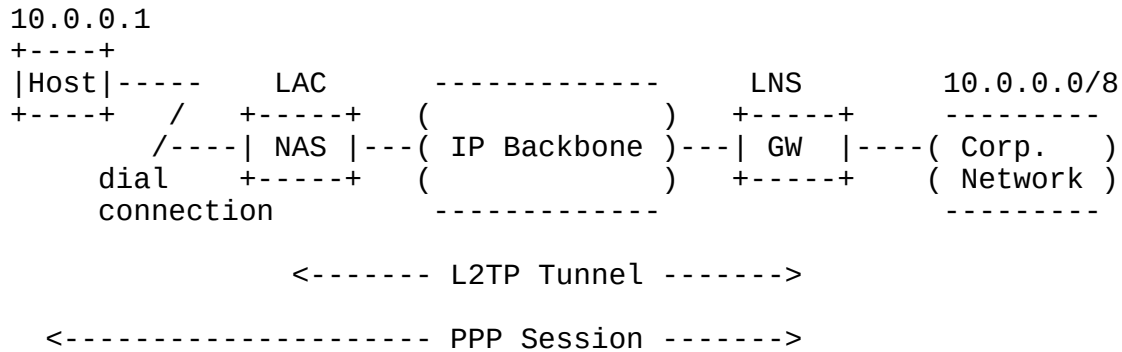


Figure 6.1: Compulsory Tunneling Example

Compulsory tunneling was originally intended for deployment on network access servers supporting wholesale dial services, allowing for remote dial access through common facilities to an enterprise site, while precluding the need for the enterprise to deploy its own dial servers. Another example of this is where an ISP outsources its own dial connectivity to an access network provider (such as a Local Exchange Carrier (LEC) in the USA) removing the need for an ISP to maintain its own dial servers and allowing the LEC to serve multiple ISPs. More recently, compulsory tunneling mechanisms have also been proposed for evolving Digital Subscriber Line (DSL) services [56], [57], which also seek to leverage the existing AAA infrastructure.

Call routing for compulsory tunnels requires that some aspect of the initial PPP call set up can be used to allow the LAC to determine the identity of the LNS. As noted in [50], these aspects can include the user identity, as determined through some aspect of the access network, including calling party number, or some attribute of the called party, such as the Fully Qualified Domain Name (FQDN) of the identity claimed during PPP authentication.

It is also possible to chain two L2TP tunnels together, whereby a LAC initiates a tunnel to an intermediate relay device, which acts as an LNS to this first LAC, and acts as a LAC to the final LNS. This may be needed in some cases due to administrative, organizational or regulatory issues pertaining to the split between access network provider, IP backbone provider and enterprise customer.

6.3 Voluntary Tunnels

Voluntary tunneling refers to the case where an individual host connects to a remote site using a tunnel originating on the host, with no involvement from intermediate network nodes, as illustrated below. The PPTP specification, parts of which have been incorporated into L2TP, was based upon a voluntary tunneling model.

As with compulsory tunneling there are different deployment scenarios possible. The diagram below shows a subscriber host accessing a corporate network with either L2TP or IPsec being used as the voluntary tunneling mechanism. Another scenario is where voluntary tunneling is used to provide a subscriber with access to a VPRN.

6.3.1 Issues with Use of L2TP for Voluntary Tunnels

The L2TP specification has support for voluntary tunneling, insofar as the LAC can be located on a host, not only on a network node. Note that such a host has two IP addresses - one for the LAC-LNS IP tunnel, and another, typically allocated via PPP, for the network to which the host is connecting. The benefits of using L2TP for voluntary tunneling are that the existing authentication and address assignment mechanisms used by PPP can be reused without modification. For example an LNS could also include a Radius client, and communicate with a Radius server to authenticate a PPP PAP or CHAP exchange, and to retrieve configuration information for the host such as its IP address and a list of DNS servers to use. This information can then be passed to the host via the PPP IPCP protocol.

10.0.0.1

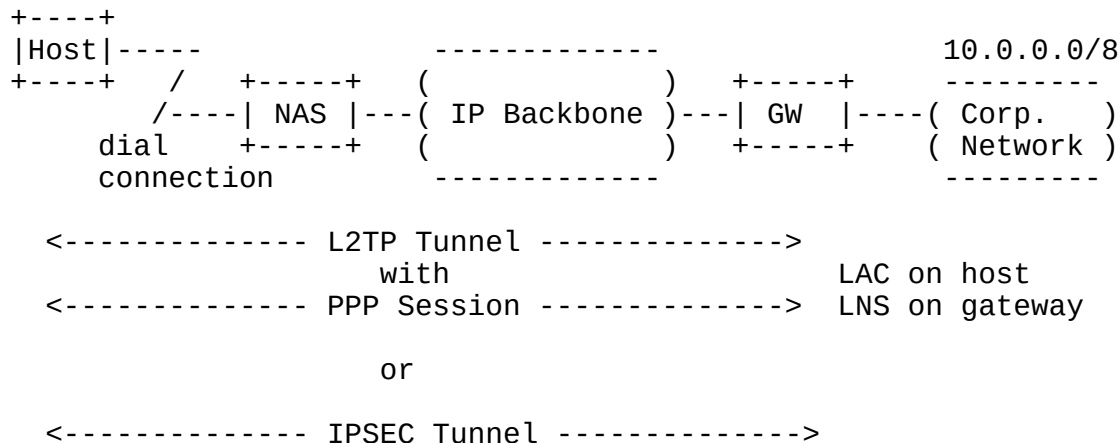


Figure 6.2: Voluntary Tunneling Example

The above procedure is not without its costs, however. There is considerable overhead with such a protocol stack, particularly when IPsec is also needed for security purposes, and given that the host may be connected via a low-bandwidth dial up link. The overhead consists of both extra headers in the data plane and extra control protocols needed in the control plane. Using L2TP for voluntary tunneling, secured with IPsec, means a web application, for example, would run over the following stack

HTTP/TCP/IP/PPP/L2TP/UDP/ESP/IP/PPP/AHDLC

It is proposed in [58] that IPsec alone be used for voluntary tunnels reducing overhead, using the following stack.

HTTP/TCP/IP/ESP/IP/PPP/AHDLC

In this case IPsec is used in tunnel mode, with the tunnel terminating either on an IPsec edge device at the enterprise site, or on the provider edge router connected to the enterprise site. There are two possibilities for the IP addressing of the host. Two IP addresses could be used, in a similar manner to the L2TP case. Alternatively the host can use a single public IP address as the source IP address in both inner and outer IP headers, with the gateway performing Network Address Translation (NAT) before forwarding the traffic to the enterprise network. To other hosts in the enterprise network the host appears to have an 'internal' IP address. Using NAT has some limitations and restrictions, also pointed out in [58].

Another area of potential problems with PPP is due to the fact that the characteristics of a link layer implemented via an L2TP tunnel over an IP backbone are quite different to a link layer run over a serial line, as discussed in the L2TP specification itself. For example, poorly chosen PPP parameters may lead to frequent resets and timeouts, particularly if compression is in use. This is because an L2TP tunnel may misorder packets, and may silently drop packets, neither of which normally occurs on serial lines. The general packet loss rate could also be significantly higher due to network congestion. Using the sequence number field in an L2TP header addresses the misordering issue, and for cases where the LAC and LNS are coincident with the PPP endpoints, as in voluntary tunneling, the sequence number field can also be used to detect a dropped packet, and to pass a suitable indication to any compression entity in use, which typically requires such knowledge in order to keep the compression histories in synchronization at both ends. (In fact this is more of an issue with compulsory tunneling since the LAC may have to deliberately issue a corrupted frame to the PPP host, to give an indication of packet loss, and some hardware may not allow this).

6.3.2 Issues with Use of IPSec for Voluntary Tunnels

If IPSec is used for voluntary tunneling, the functions of user authentication and host configuration, achieved by means of PPP when using L2TP, still need to be carried out. A distinction needs to be drawn here between machine authentication and user authentication. 'Two factor' authentication is carried out on the basis of both something the user has, such as a machine or smartcard with a digital certificate, and something the user knows, such as a password. (Another example is getting money from an bank ATM machine - you need a card and a PIN number). Many of the existing legacy schemes currently in use to perform user authentication are asymmetric in nature, and are not supported by IKE. For remote access the most common existing user authentication mechanism is to use PPP between the user and access server, and Radius between the access server and authentication server. The authentication exchanges that occur in this case, e.g. a PAP or CHAP exchange, are asymmetric. Also CHAP supports the ability for the network to reauthenticate the user at any time after the initial session has been established, to ensure that the current user is the same person that initiated the session.

While IKE provides strong support for machine authentication, it has only limited support for any form of user authentication and has no support for asymmetric user authentication. While a user password can be used to derive a key used as a preshared key, this cannot be used with IKE Main Mode in a remote access environment, as the user will not have a fixed IP address, and while Aggressive Mode can be used instead, this affords no identity protection. To this end there have been a number of proposals to allow for support of legacy asymmetric user level authentication schemes with IPSec. [59] defines a new IKE message exchange - the transaction exchange - which allows for both Request/Reply and Set/Acknowledge message sequences, and it also defines attributes that can be used for client IP stack configuration. [60] and [61] describe mechanisms that use the transaction message exchange, or a series of such exchanges, carried out between the IKE Phase 1 and Phase 2 exchanges, to perform user authentication. A different approach, that does not extend the IKE protocol itself, is described in [62]. With this approach a user establishes a Phase 1 SA with a security gateway and then sets up a Phase 2 SA to the gateway, over which an existing authentication protocol is run. The gateway acts as a proxy and relays the protocol messages to an authentication server.

In addition there have also been proposals to allow the remote host to be configured with an IP address and other configuration information over IPSec. For example [63] describes a method whereby a remote host first establishes a Phase 1 SA with a security gateway and then sets up a Phase 2 SA to the gateway, over which the DHCP

protocol is run. The gateway acts as a proxy and relays the protocol messages to the DHCP server. Again, like [62], this proposal does not involve extensions to the IKE protocol itself.

Another aspect of PPP functionality that may need to be supported is multiprotocol operation, as there may be a need to carry network layer protocols other than IP, and even to carry link layer protocols (e.g. ethernet) as would be needed to support bridging over IPSec. This is discussed in section 3.1.4.

The methods of supporting legacy user authentication and host configuration capabilities in a remote access environment are currently being discussed in the IPSec working group.

6.4 Networked Host Support

The current PPP based dial model assumes a host directly connected to a connection oriented dial access network. Recent work on new access technologies such as DSL have attempted to replicate this model [57], so as to allow for the re-use of existing AAA systems. The proliferation of personal computers, printers and other network appliances in homes and small businesses, and the ever lowering costs of networks, however, are increasingly challenging the directly connected host model. Increasingly, most hosts will access the Internet through small, typically Ethernet, local area networks.

There is hence interest in means of accommodating the existing AAA infrastructure within service providers, whilst also supporting multiple networked hosts at each customer site. The principal complication with this scenario is the need to support the login dialogue, through which the appropriate AAA information is exchanged. A number of proposals have been made to address this scenario:

6.4.1 Extension of PPP to Hosts Through L2TP

A number of proposals (e.g. [56]) have been made to extend L2TP over Ethernet so that PPP sessions can run from networked hosts out to the network, in much the same manner as a directly attached host.

6.4.2 Extension of PPP Directly to Hosts:

There is also a specification for mapping PPP directly onto Ethernet (PPPOE) [64] which uses a broadcast mechanism to allow hosts to find appropriate access servers with which to connect. Such servers could then further tunnel, if needed, the PPP sessions using L2TP or a similar mechanism.

6.4.3 Use of IPSec

The IPSec based voluntary tunneling mechanisms discussed above can be used either with networked or directly connected hosts.

Note that all of these methods require additional host software to be used, which implements either LAC, PPPoE client or IPSec client functionality.

6.5 Recommendations

The L2TP specification has been finalized and will be widely used for compulsory tunneling. As discussed in section 3.2, defining specific modes of operation for IPSec when used to secure L2TP would be beneficial.

Also, for voluntary tunneling using IPSec, completing the work needed to provide support for the following areas would be useful

- asymmetric / legacy user authentication (6.3)
- host address assignment and configuration (6.3)

along with any other issues specifically related to the support of remote hosts. Currently as there are many different non-interoperable proprietary solutions in this area.

7.0 VPN Types: Virtual Private LAN Segment

A Virtual Private LAN Segment (VPLS) is the emulation of a LAN segment using Internet facilities. A VPLS can be used to provide what is sometimes known also as a Transparent LAN Service (TLS), which can be used to interconnect multiple stub CPE nodes, either bridges or routers, in a protocol transparent manner. A VPLS emulates a LAN segment over IP, in the same way as protocols such as LANE emulate a LAN segment over ATM. The primary benefits of a VPLS are complete protocol transparency, which may be important both for multiprotocol transport and for regulatory reasons in particular service provider contexts.

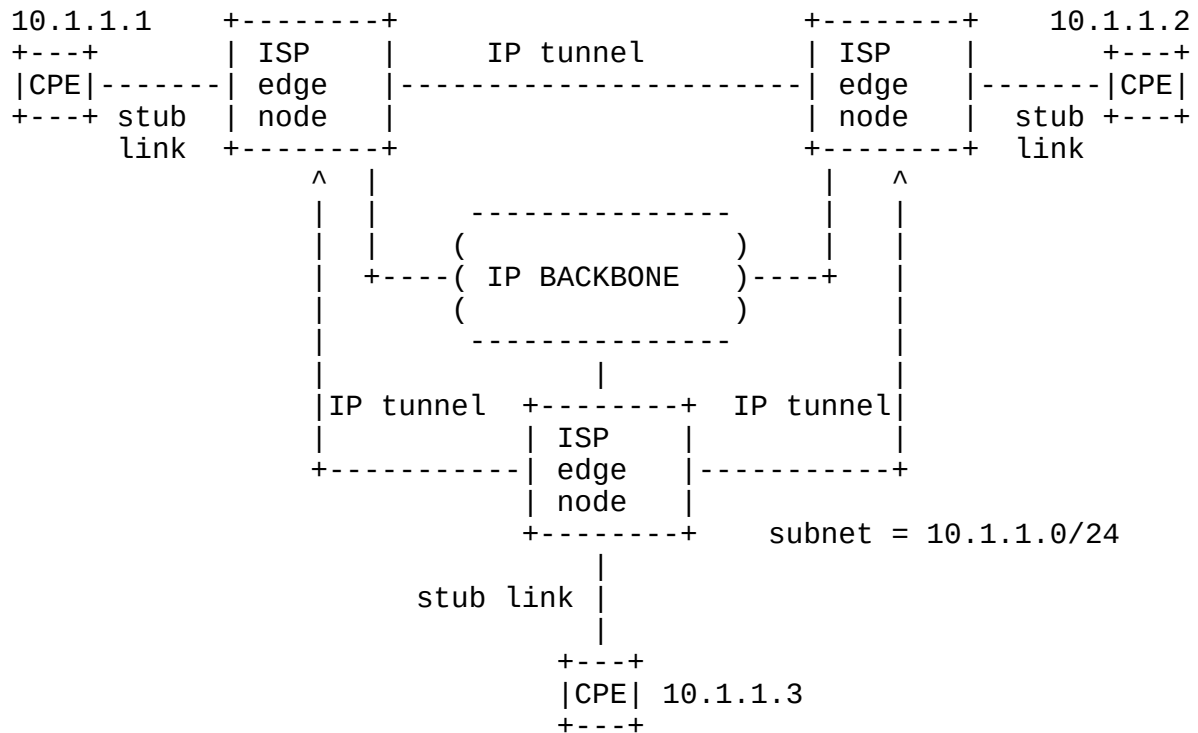


Figure 7.1: VPLS Example

7.1 VPLS Requirements

Topologically and operationally a VPLS can be most easily modeled as being essentially equivalent to a VPRN, except that each VPLS edge node implements link layer bridging rather than network layer forwarding. As such, most of the VPRN tunneling and configuration mechanisms discussed previously can also be used for a VPLS, with the appropriate changes to accommodate link layer, rather than network layer, packets and addressing information. The following sections discuss the primary changes needed in VPRN operation to support VPLSs.

7.1.1 Tunneling Protocols

The tunneling protocols employed within a VPLS can be exactly the same as those used within a VPRN, if the tunneling protocol permits the transport of multiprotocol traffic, and this is assumed below.

7.1.2 Multicast and Broadcast Support

A VPLS needs to have a broadcast capability. This is needed both for broadcast frames, and for link layer packet flooding, where a unicast frame is flooded because the path to the destination link layer address is unknown. The address resolution protocols that run over a bridged network typically use broadcast frames (e.g. ARP). The same set of possible multicast tunneling mechanisms discussed earlier for VPRNs apply also to a VPLS, though the generally more frequent use of broadcast in VPLSs may increase the pressure for native multicast support that reduces, for instance, the burden of replication on VPLS edge nodes.

7.1.3 VPLS Membership Configuration and Topology

The configuration of VPLS membership is analogous to that of VPRNs since this generally requires only knowledge of the local VPN link assignments at any given VPLS edge node, and the identity of, or route to, the other edge nodes in the VPLS; in particular, such configuration is independent of the nature of the forwarding at each VPN edge node. As such, any of the mechanisms for VPN member configuration and dissemination discussed for VPRN configuration can also be applied to VPLS configuration. Also as with VPRNs, the topology of the VPLS could be easily manipulated by controlling the configuration of peer nodes at each VPLS edge node, assuming that the membership dissemination mechanism was such as to permit this. It is likely that typical VPLSs will be fully meshed, however, in order to preclude the need for traffic between two VPLS nodes to transit through another VPLS node, which would then require the use of the Spanning Tree protocol [65] for loop prevention.

7.1.4 CPE Stub Node Types

A VPLS can support either bridges or routers as a CPE device.

CPE routers would peer transparently across a VPLS with each other without requiring any router peering with any nodes within the VPLS. The same scalability issues that apply to a full mesh topology for VPRNs, apply also in this case, only that now the number of peering routers is potentially greater, since the ISP edge device is no longer acting as an aggregation point.

With CPE bridge devices the broadcast domain encompasses all the CPE sites as well as the VPLS itself. There are significant scalability constraints in this case, due to the need for packet flooding, and

the fact that any topology change in the bridged domain is not localized, but is visible throughout the domain. As such this scenario is generally only suited for support of non-routable protocols.

The nature of the CPE impacts the nature of the encapsulation, addressing, forwarding and reachability protocols within the VPLS, and are discussed separately below.

7.1.5 Stub Link Packet Encapsulation

7.1.5.1 Bridge CPE

In this case, packets sent to and from the VPLS across stub links are link layer frames, with a suitable access link encapsulation. The most common case is likely to be ethernet frames, using an encapsulation appropriate to the particular access technology, such as ATM, connecting the CPE bridges to the VPLS edge nodes. Such frames are then forwarded at layer 2 onto a tunnel used in the VPLS. As noted previously, this does mandate the use of an IP tunneling protocol which can transport such link layer frames. Note that this does not necessarily mandate, however, the use of a protocol identification field in each tunnel packet, since the nature of the encapsulated traffic (e.g. ethernet frames) could be indicated at tunnel setup.

7.1.5.2 Router CPE

In this case, typically, CPE routers send link layer packets to and from the VPLS across stub links, destined to the link layer addresses of their peer CPE routers. Other types of encapsulations may also prove feasible in such a case, however, since the relatively constrained addressing space needed for a VPLS to which only router CPE are connected, could allow for alternative encapsulations, as discussed further below.

7.1.6 CPE Addressing and Address Resolution

7.1.6.1 Bridge CPE

Since a VPLS operates at the link layer, all hosts within all stub sites, in the case of bridge CPE, will typically be in the same network layer subnet. (Multinetting, whereby multiple subnets operate over the same LAN segment, is possible, but much less common). Frames are forwarded across and within the VPLS based upon the link layer addresses - e.g. IEEE MAC addresses - associated with the individual hosts. The VPLS needs to support broadcast traffic, such as that typically used for the address resolution mechanism used

to map the host network addresses to their respective link addresses. The VPLS forwarding and reachability algorithms also need to be able to accommodate flooded traffic.

7.1.6.2 Router CPE

A single network layer subnet is generally used to interconnect router CPE devices, across a VPLS. Behind each CPE router are hosts in different network layer subnets. CPE routers transfer packets across the VPLS by mapping next hop network layer addresses to the link layer addresses of a router peer. A link layer encapsulation is used, most commonly ethernet, as for the bridge case.

As noted above, however, in cases where all of the CPE nodes connected to the VPLS are routers, then it may be possible, due to the constrained addressing space of the VPLS, to use encapsulations that use a different address space than normal MAC addressing. See, for instance, [11], for a proposed mechanism for VPLSs over MPLS networks, leveraging earlier work on VPRN support over MPLS [38], which proposes MPLS as the tunneling mechanism, and locally assigned MPLS labels as the link layer addressing scheme to identify the CPE LSR routers connected to the VPLS.

7.1.7 VPLS Edge Node Forwarding and Reachability Mechanisms

7.1.7.1 Bridge CPE

The only practical VPLS edge node forwarding mechanism in this case is likely to be standard link layer packet flooding and MAC address learning, as per [65]. As such, no explicit intra-VPLS reachability protocol will be needed, though there will be a need for broadcast mechanisms to flood traffic, as discussed above. In general, it may not prove necessary to also implement the Spanning Tree protocol between VPLS edge nodes, if the VPLS topology is such that no VPLS edge node is used for transit traffic between any other VPLS edge nodes - in other words, where there is both full mesh connectivity and transit is explicitly precluded. On the other hand, the CPE bridges may well implement the spanning tree protocol in order to safeguard against 'backdoor' paths that bypass connectivity through the VPLS.

7.1.7.2 Router CPE

Standard bridging techniques can also be used in this case. In addition, the smaller link layer address space of such a VPLS may also permit other techniques, with explicit link layer routes between CPE routers. [11], for instance, proposes that MPLS LSPs be set up, at the insertion of any new CPE router into the VPLS, between all CPE

LSRs. This then precludes the need for packet flooding. In the more general case, if stub link reachability mechanisms were used to configure VPLS edge nodes with the link layer addresses of the CPE routers connected to them, then modifications of any of the intra-VPN reachability mechanisms discussed for VPRNs could be used to propagate this information to each other VPLS edge node. This would then allow for packet forwarding across the VPLS without flooding.

Mechanisms could also be developed to further propagate the link layer addresses of peer CPE routers and their corresponding network layer addresses across the stub links to the CPE routers, where such information could be inserted into the CPE router's address resolution tables. This would then also preclude the need for broadcast address resolution protocols across the VPLS.

Clearly there would be no need for the support of spanning tree protocols if explicit link layer routes were determined across the VPLS. If normal flooding mechanisms were used then spanning tree would only be required if full mesh connectivity was not available and hence VPLS nodes had to carry transit traffic.

7.2 Recommendations

There is significant commonality between VPRNs and VPLSs, and, where possible, this similarity should be exploited in order to reduce development and configuration complexity. In particular, VPLSs should utilize the same tunneling and membership configuration mechanisms, with changes only to reflect the specific characteristics of VPLSs.

8.0 Summary of Recommendations

In this document different types of VPNs have been discussed individually, but there are many common requirements and mechanisms that apply to all types of VPNs, and many networks will contain a mix of different types of VPNs. It is useful to have as much commonality as possible across these different VPN types. In particular, by standardizing a relatively small number of mechanisms, it is possible to allow a wide variety of VPNs to be implemented.

The benefits of adding support for the following mechanisms should be carefully examined.

For IKE/IPSec:

- the transport of a VPN-ID when establishing an SA (3.1.2)
- a null encryption and null authentication option (3.1.3)

- multiprotocol operation (3.1.4)
- frame sequencing (3.1.5)
- asymmetric / legacy user authentication (6.3)
- host address assignment and configuration (6.3)

For L2TP:

- defining modes of operation of IPSec when used to support L2TP (3.2)

For VPNs generally:

- defining a VPN membership information configuration and dissemination mechanism, that uses some form of directory or MIB (5.3.2)
- ensure that solutions developed, as far as possible, are applicable to different types of VPNs, rather than being specific to a single type of VPN.

9.0 Security Considerations

Security considerations are an integral part of any VPN mechanisms, and these are discussed in the sections describing those mechanisms.

10.0 Acknowledgements

Thanks to Anthony Alles, of Nortel Networks, for his invaluable assistance with the generation of this document, and who developed much of the material on which early versions of this document were based. Thanks also to Joel Halpern for his helpful review comments.

11.0 References

- [1] ATM Forum. "LAN Emulation over ATM 1.0", af-lane-0021.000, January 1995.
- [2] ATM Forum. "Multi-Protocol Over ATM Specification v1.0", af-mpoa-0087.000, June 1997.
- [3] Ferguson, P. and Huston, G. "What is a VPN?", Revision 1, April 1 1998; <http://www.employees.org/~ferguson/vpn.pdf>.

- [4] Rekhter, Y., Moskowitz, B., Karrenberg, D., de Groot, G. and E. Lear, "Address Allocation for Private Internets", BCP 5, RFC 1918, February 1996.
- [5] Kent, S. and R. Atkinson, "Security Architecture for the Internet Protocol", RFC 2401, November 1998.
- [6] Perkins, C., "IP Encapsulation within IP", RFC 2003, October 1996.
- [7] Hanks, S., Li, T., Farinacci, D. and P. Traina, "Generic Routing Encapsulation (GRE)", RFC 1701, October 1994.
- [8] Townsley, W., Valencia, A., Rubens, A., Pall, G., Zorn, G. and B. Palter, "Layer Two Tunneling Protocol "L2TP"", RFC 2661, August 1999.
- [9] Rosen, E., et al., "Multiprotocol Label Switching Architecture", Work in Progress.
- [10] Heinanen, J., et al., "MPLS Mappings of Generic VPN Mechanisms", Work in Progress.
- [11] Jamieson, D., et al., "MPLS VPN Architecture", Work in Progress.
- [12] Casey, L., et al., "IP VPN Realization using MPLS Tunnels", Work in Progress.
- [13] Li, T. "CPE based VPNs using MPLS", Work in Progress.
- [14] Muthukrishnan, K. and A. Malis, "Core MPLS IP VPN Architecture", Work in Progress.
- [15] Rosen, E. and Y. Rekhter, "BGP/MPLS VPNs", RFC 2547, March 1999.
- [16] Fox, B. and B. Gleeson, "Virtual Private Networks Identifier", RFC 2685, September 1999.
- [17] Petri, B. (editor) "MPOA v1.1 Addendum on VPN support", ATM Forum, af-mpoa-0129.000.
- [18] Harkins, D. and C. Carrel, "The Internet Key Exchange (IKE)", RFC 2409, November 1998.
- [19] Calhoun, P., et al., "Tunnel Establishment Protocol", Work in Progress.

- [20] Andersson, L., et al., "LDP Specification", Work in Progress.
- [21] Jamoussi, B., et al., "Constraint-Based LSP Setup using LDP" Work in Progress.
- [22] Awduche, D., et al., "Extensions to RSVP for LSP Tunnels", Work in Progress.
- [23] Kent, S. and R. Atkinson, "IP Encapsulating Security Protocol (ESP)", RFC 2406, November 1998.
- [24] Simpson, W., Editor, "The Point-to-Point Protocol (PPP)", STD 51, RFC 1661, July 1994.
- [25] Perez, M., Liaw, F., Mankin, A., Hoffman, E., Grossman, D. and A. Malis, "ATM Signalling Support for IP over ATM", RFC 1755, February 1995.
- [26] Malkin, G. "RIP Version 2 Carrying Additional Information", RFC 1723, November 1994.
- [27] Moy, J., "OSPF Version 2", STD 54, RFC 2328, April 1998.
- [28] Shacham, A., Monsour, R., Pereira, R. and M. Thomas, "IP Payload Compression Protocol (IPComp)", RFC 2393, December 1998.
- [29] Duffield N., et al., "A Performance Oriented Service Interface for Virtual Private Networks", Work in Progress.
- [30] Jacobson, V., Nichols, K. and B. Poduri, "An Expedited Forwarding PHB", RFC 2598, June 1999.
- [31] Casey, L., "An extended IP VPN Architecture", Work in Progress.
- [32] Rekhter, Y., and T. Li, "A Border Gateway Protocol 4 (BGP-4)", RFC 1771, March 1995.
- [33] Grossman, D. and J. Heinanen, "Multiprotocol Encapsulation over ATM Adaptation Layer 5", RFC 2684, September 1999.
- [34] Wahl, M., Howes, T. and S. Kille, "Lightweight Directory Access Protocol (v3)", RFC 2251, December 1997.
- [35] Boyle, J., et al., "The COPS (Common Open Policy Service) Protocol", RFC 2748, January 2000.
- [36] MacRae, M. and S. Ayandeh, "Using COPS for VPN Connectivity" Work in Progress.

- [37] Droms, R., "Dynamic Host Configuration Protocol", RFC 2131, March 1997.
- [38] Heinanen, J. and E. Rosen, "VPN Support with MPLS", Work in Progress.
- [39] Estrin, D., Farinacci, D., Helmy, A., Thaler, D., Deering, S., Handley, M., Jacobson, V., Liu, C., Sharma, P. and L. Wei, "Protocol Independent Multicast-Sparse Mode (PIM-SM): Protocol Specification", RFC 2362, June 1998.
- [40] Waitzman, D., Partridge, C., and S. Deering, "Distance Vector Multicast Routing Protocol", RFC 1075, November 1988.
- [41] Fenner, W., "IGMP-based Multicast Forwarding (IGMP Proxying)", Work in Progress.
- [42] Wallner, D., Harder, E. and R. Agee, "Key Management for Multicast: Issues and Architectures", RFC 2627, June 1999.
- [43] Hardjono, T., et al., "Secure IP Multicast: Problem areas, Framework, and Building Blocks", Work in Progress.
- [44] Rigney, C., Rubens, A., Simpson, W. and S. Willens, "Remote Authentication Dial In User Service (RADIUS)", RFC 2138, April 1997.
- [45] Valencia, A., Littlewood, M. and T. Kolar, "Cisco Layer Two Forwarding (Protocol) "L2F"", RFC 2341, May 1998.
- [46] Hamzeh, K., Pall, G., Verthein, W., Taarud, J., Little, W. and G. Zorn, "Point-to-Point Tunneling Protocol (PPTP)", RFC 2637, July 1999.
- [47] Patel, B., et al., "Securing L2TP using IPSEC", Work in Progress.
- [48] Srisuresh, P., "Secure Remote Access with L2TP", Work in Progress.
- [49] Calhoun, P., et al., "Layer Two Tunneling Protocol "L2TP" Security Extensions for Non-IP networks", Work in Progress.
- [50] Aboba, B. and Zorn, G. "Implementation of PPTP/L2TP Compulsory Tunneling via RADIUS", Work in progress.
- [51] Aboba, B. and G. Zorn, "Criteria for Evaluating Roaming Protocols", RFC 2477, January 1999.

- [52] Shea, R., "L2TP-over-IP Path MTU Discovery (L2TPMTU)", Work in Progress.
- [53] Sklower, K., Lloyd, B., McGregor, G., Carr, D. and T. Coradetti, "The PPP Multilink Protocol (MP)", RFC 1990, August 1996.
- [54] Richards, C. and K. Smith, "The PPP Bandwidth Allocation Protocol (BAP) The PPP Bandwidth Allocation Control Protocol (BACP)", RFC 2125, March 1997.
- [55] Calhoun, P. and K. Peirce, "Layer Two Tunneling Protocol "L2TP" IP Differential Services Extension", Work in Progress.
- [56] ADSL Forum. "An Interoperable End-to-end Broadband Service Architecture over ADSL Systems (Version 3.0)", ADSL Forum 97-215.
- [57] ADSL Forum. "Core Network Architectures for ADSL Access Systems (Version 1.01)", ADSL Forum 98-017.
- [58] Gupta, V., "Secure, Remote Access over the Internet using IPsec", Work in Progress.
- [59] Pereira, R., et al., "The ISAKMP Configuration Method", Work in Progress.
- [60] Pereira, R. and S. Beaulieu, "Extended Authentication Within ISAKMP/Oakley", Work in Progress.
- [61] Litvin, M., et al., "A Hybrid Authentication Mode for IKE", Work in Progress.
- [62] Kelly, S., et al., "User-level Authentication Mechanisms for IPsec", Work in Progress.
- [63] Patel, B., et al., "DHCP Configuration of IPSEC Tunnel Mode", Work in Progress.
- [64] Mamakos, L., Lidl, K., Evarts, J., Carrel, D., Simone, D. and R. Wheeler, "A Method for Transmitting PPP Over Ethernet (PPPoE)", RFC 2516, February 1999.
- [65] ANSI/IEEE - 10038: 1993 (ISO/IEC) Information technology - Telecommunications and information exchange between systems - Local area networks - Media access control (MAC) bridges, ANSI/IEEE Std 802.1D, 1993 Edition.

12.0 Author Information

Bryan Gleeson
Nortel Networks
4500 Great America Parkway
Santa Clara CA 95054
USA

Phone: +1 (408) 548 3711
EMail: bgleeson@shastanets.com

Juha Heinanen
Telia Finland, Inc.
Myrmaentie 2
01600 VANTAA
Finland

Phone: +358 303 944 808
EMail: jh@telia.fi

Arthur Lin
Nortel Networks
4500 Great America Parkway
Santa Clara CA 95054
USA

Phone: +1 (408) 548 3788
EMail: alin@shastanets.com

Grenville Armitage
Bell Labs Research Silicon Valley
Lucent Technologies
3180 Porter Drive,
Palo Alto, CA 94304
USA

EMail: gja@lucent.com

Andrew G. Malis
Lucent Technologies
1 Robbins Road
Westford, MA 01886
USA

Phone: +1 978 952 7414
EMail: amalis@lucent.com

13.0 Full Copyright Statement

Copyright (C) The Internet Society (2000). All Rights Reserved.

This document and translations of it may be copied and furnished to others, and derivative works that comment on or otherwise explain it or assist in its implementation may be prepared, copied, published and distributed, in whole or in part, without restriction of any kind, provided that the above copyright notice and this paragraph are included on all such copies and derivative works. However, this document itself may not be modified in any way, such as by removing the copyright notice or references to the Internet Society or other Internet organizations, except as needed for the purpose of developing Internet standards in which case the procedures for copyrights defined in the Internet Standards process must be followed, or as required to translate it into languages other than English.

The limited permissions granted above are perpetual and will not be revoked by the Internet Society or its successors or assigns.

This document and the information contained herein is provided on an "AS IS" basis and THE INTERNET SOCIETY AND THE INTERNET ENGINEERING TASK FORCE DISCLAIMS 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.

Acknowledgement

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

