

DMVPN Design Overview

This chapter provides an overview of the DMVPN design topology and characteristics. [Chapter 2, “DMVPN Design and Implementation,”](#) provides more detail on the design considerations. [Chapter 3, “Scalability Considerations,”](#) then presents Cisco product options for deploying the design.

Overview

The primary topology discussed is a hub-and-spoke deployment model in which the primary enterprise resources are located in a large central site, with a number of smaller sites or branch offices connected directly to the central site over a VPN. However, in some scenarios, a spoke-to-spoke deployment model can be used, which provides the ability to create temporary connections between branch sites directly using IPsec encryption. Both DMVPN deployment models are shown in [Figure 1-1](#).

Figure 1-1 *DMVPN Deployment Models*



Starting Assumptions

The design approach presented in this design guide makes the following starting assumptions:

- The design supports a typical converged traffic profile for customers (see [Chapter 4, “Scalability Test Results \(Unicast Only\).”](#))
- The customer has a need for diverse traffic requirements such as IP multicast and support for routing. The use of mGRE and a routing protocol are also discussed in more detail in [Chapter 2, “DMVPN Design and Implementation.”](#)
- Cisco products should be maintained at reasonable CPU utilization levels. This is discussed in more detail in [Chapter 3, “Scalability Considerations,”](#) including recommendations for both headend and branch routers, and software revisions.
- Although costs were certainly considered, the design recommendations assume that the customer deploys current VPN technologies, including hardware-accelerated encryption.
- Voice over IP (VoIP) and video are assumed to be requirements in the network. Detailed design considerations for handling VoIP and other latency sensitive traffic are not explicitly addressed in this design guide, but may be found in *Voice and Video Enabled IPsec VPN (V3PN) Design Guide* at the following URL:
http://www.cisco.com/en/US/docs/solutions/Enterprise/WAN_and_MAN/V3PN_SRND/V3PN_SRND.html.
- This design is targeted for deployment by enterprise-owned VPNs; however, the concepts and conclusions are valid regardless of the ownership of the edge tunneling equipment, and are therefore valuable for service provider-managed VPNs as well.

Design Components

VPNs provide an alternate to traditional WAN technologies such as leased lines, Frame Relay, and ATM. VPN technology allows private WANs to exist over a public transport such as the Internet. LAN-to-LAN VPNs are primarily deployed to connect branch office locations to the central site (or sites) of an enterprise.

The requirements of enterprise customers for traditional private WAN services such as multiprotocol support, high availability, scalability, and security are also requirements for VPNs. VPNs can often meet these requirements more cost-effectively and with greater flexibility than private WAN services.

The following are key components of this DMVPN design:

- Cisco high-end VPN routers serving as VPN headend termination devices at a central campus (headend devices)
- Cisco VPN access routers serving as VPN branch termination devices at branch office locations (branch devices)
- DMVPN hub-and-spoke to perform headend-to-branch interconnections
- DMVPN spoke-to-spoke to perform branch-to-branch interconnections (optional)
- Internet services procured from a third-party ISP (or ISPs) serving as the WAN interconnection medium

Cisco VPN routers are a good choice for VPN deployments because they can accommodate any network requirement traditionally provided by a Frame Relay or private line network. These requirements include support for multicast, latency-sensitive traffic, and routing protocols. See [Chapter 3, “Scalability Considerations,”](#) for a discussion on selection of headend and branch products.

Design Topologies

In a DMVPN design, the following two topologies can be implemented:

- Dual hub-dual DMVPN cloud
- Dual hub-single DMVPN cloud

**Note**

Wherever dual hub is mentioned it means two or more hubs, there is not a limitation of only two hubs in either the dual hub-dual DMVPN or dual hub-single DMVPN cloud topologies. In this document only two hubs are used for simplicity of the functional explanation.

In both topologies at least two hubs or headends are recommended for redundancy. A DMVPN cloud is a collection of routers that is configured either with a multipoint GRE (mGRE) interface or point-to-point (p2p) GRE interface (or combination of the two) that share the same IP address subnet. High availability is provided through the use of a second hub router, which may be on the same DMVPN subnet as the primary hub router. This is commonly referred to as a single DMVPN cloud topology. The second hub router can also service its own DMVPN subnet, which is known as a dual DMVPN cloud topology. A dual hub-single DMVPN and dual hub-dual DMVPN cloud topologies rely on routing protocols running inside of the tunnels to determine tunnel path selection.

A DMVPN cloud topology can support either a hub-and-spoke or spoke-to-spoke deployment model. In a hub-and-spoke deployment model, each headend uses an mGRE interface and each branch uses a p2pGRE or mGRE interface. In a spoke-to-spoke deployment model, both the headend and the branch must use mGRE interfaces.

[Figure 1-2](#) and [Figure 1-3](#) show the two DMVPN cloud topologies. More details on the various deployment models under this topology is discussed in the next section.

Figure 1-2 **Dual DMVPN Cloud Topology**



Figure 1-3 *Single DMVPN Cloud Topology*



The difference between the two topologies is most apparent on the branch router. With a single DMVPN subnet, the branch router has a single mGRE tunnel, and both headends are mapped to this tunnel through this mGRE interface. In a dual DMVPN topology, the branch router has a unique tunnel pointing to a unique headend. Standard routing protocols such as EIGRP, BGP or OSPF are used to determine the active hub over either topology. In the single DMVPN topology the hubs will appear as two different next-hops via the one mGRE tunnel interface, whereas in the dual DMVPN topology the hubs will appear as two different next-hops via the two p-pGRE or mGRE interfaces. Note, both hubs may not appear in the routing table depending on how the routing protocol metrics for the hubs are configured.

In general the single DMVPN cloud topology is best when dynamic spoke-spoke tunnels are required, because spoke-spoke tunnels can only be built within a DMVPN cloud not between DMVPN clouds. The dual DMVPN cloud topology is often easier for hub-and-spoke only networks, in this case it can be easier to configure the routing protocol to prefer one DMVPN cloud (hub) over the other since the routing protocol receives routing information from the hubs on different tunnel interfaces. Though, either DMVPN cloud topology can be configured for either style network (hub-and-spoke only or spoke-spoke). These DMVPN cloud topologies can be used in combination to meet the requirements for more complex networks.

Dual DMVPN Cloud Topology

The following two deployment models can be implemented in a dual DMVPN cloud topology design:

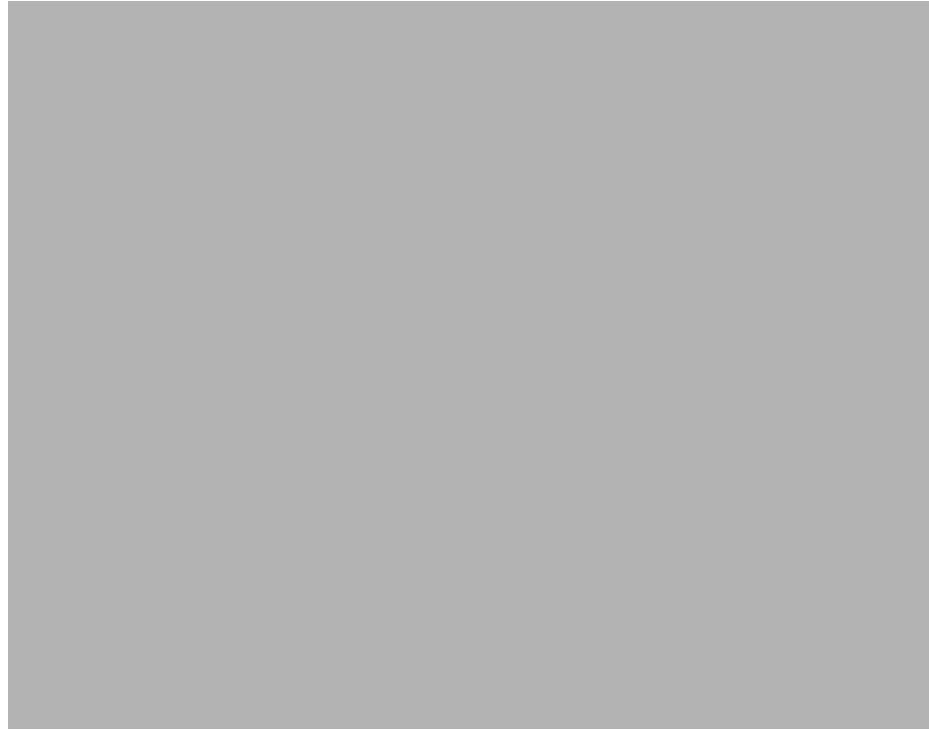
- Hub-and-spoke
- Spoke-to-spoke

Each of these deployment models is discussed in the following sections.

Dual DMVPN Cloud Topology—Hub-and-Spoke Deployment Model

A dual DMVPN cloud topology hub-and-spoke deployment model consists of two headend routers (Hub 1 and Hub 2), each with one mGRE tunnel interface(s) that connect to all branch routers (see [Figure 1-4](#)).

Figure 1-4 *Dual DMVPN Cloud Topology—Hub-and-Spoke Deployment Model*



Each DMVPN cloud represents a unique IP subnet. One DMVPN cloud is considered the primary, which all branch traffic transits. Each branch is configured with two p2p GRE tunnel interfaces, with one going to each respective headend. In this deployment model, there are no tunnels between branches. Inter-branch communication is provided through the hub routers. This closely matches traditional Frame Relay networks. Routing metrics are used to steer traffic to the primary headend router (Hub 1).

Hub-and-Spoke Deployment Model—Headend System Architectures

The following two headend system architectures can be implemented with hub-and-spoke topologies, depending on the scalability requirements:

- Single Tier Headend Architecture
- Dual Tier Headend Architecture

Single Tier Headend Architecture

In a Single Tier Headend Architecture, the mGRE and crypto functionally co-exist on the same router CPU. [Figure 1-5](#) shows this hub-and-spoke topology.

Figure 1-5 **Single Tier Headend Architecture**



In [Figure 1-5](#), the solution is a dual DMVPN cloud topology with the hub-and-spoke deployment model. Both headends are mGRE and crypto tunnel aggregation routers servicing multiple mGRE tunnels for a prescribed number of branch office locations. In addition to terminating the VPN tunnels at the central site, headends can advertise branch routes using IP routing protocols such as EIGRP or OSPF, regardless of which DMVPN cloud path selection is chosen. This is the recommended architecture.

Dual Tier Headend Architecture

In a Dual Tier Headend Architecture, the mGRE and crypto functionally does not co-exist on the same router CPU. [Figure 1-6](#) shows this hub-and-spoke topology.

Figure 1-6 Dual Tier Headend Architecture

In [Figure 1-6](#), the solution is a dual DMVPN cloud topology with the hub-and-spoke deployment model. There are separate mGRE headends and crypto headends that together service multiple mGRE tunnels for a prescribed number of branch office locations. The crypto headends terminate the VPN tunnels at the central site from each branch location and then forward the traffic to the mGRE headends that advertise branch routes using IP routing protocols such as EIGRP or OSPF. This is no longer a recommended architecture, because it breaks communication between NHRP and IPsec, allowing the possibility for black holing traffic through the tunnels for an extended time. Though, the routing protocol running through the tunnel would help to protect against this.

Dual DMVPN Cloud Topology—Hub-and-Spoke Deployment Model Branch Router Considerations

Branches in a dual DMVPN cloud topology with the hub-and-spoke deployment model use p2p GRE over IPsec tunnel(s) from the branch office locations to the central site. In addition to terminating the VPN tunnels, the branch router often provides WAN access, and in some implementations may serve as a firewall.

The public IP address of the branch router is either a statically-defined or a dynamically-assigned IP address. Both the p2p GRE and crypto tunnels are sourced from the public IP address interface. This address is registered with the headend, which provides a mapping to the branch private tunnel interface address.

Dual DMVPN Cloud Topology—Spoke-to-Spoke Deployment Model

A dual DMVPN cloud topology with the spoke-to-spoke deployment model consists of two headend routers (Hub 1 and Hub 2), each with one or more mGRE tunnel interface(s) that connect to all branch routers (see [Figure 1-7](#)). Each DMVPN cloud represents a unique IP subnet. One DMVPN cloud is considered the primary, which all branch traffic transits. On each branch router, there is an mGRE

interface into each DMVPN cloud for redundancy. All branch-to-branch communications transit through the primary headend until the dynamic spoke-to-spoke tunnel is created. The dynamic spoke-to-spoke tunnels must be within a single DMVPN cloud or subnet. Spoke-to-spoke tunnels are not possible between DMVPN clouds. Also, in some cases a two spokes site can end up building two spoke-spoke tunnels one over each DMVPN cloud. This will function, but does use more resources on the spokes routers and can complicate troubleshooting.

Figure 1-7 *Dual DMVPN Cloud Topology—Spoke-to-Spoke Deployment Model*

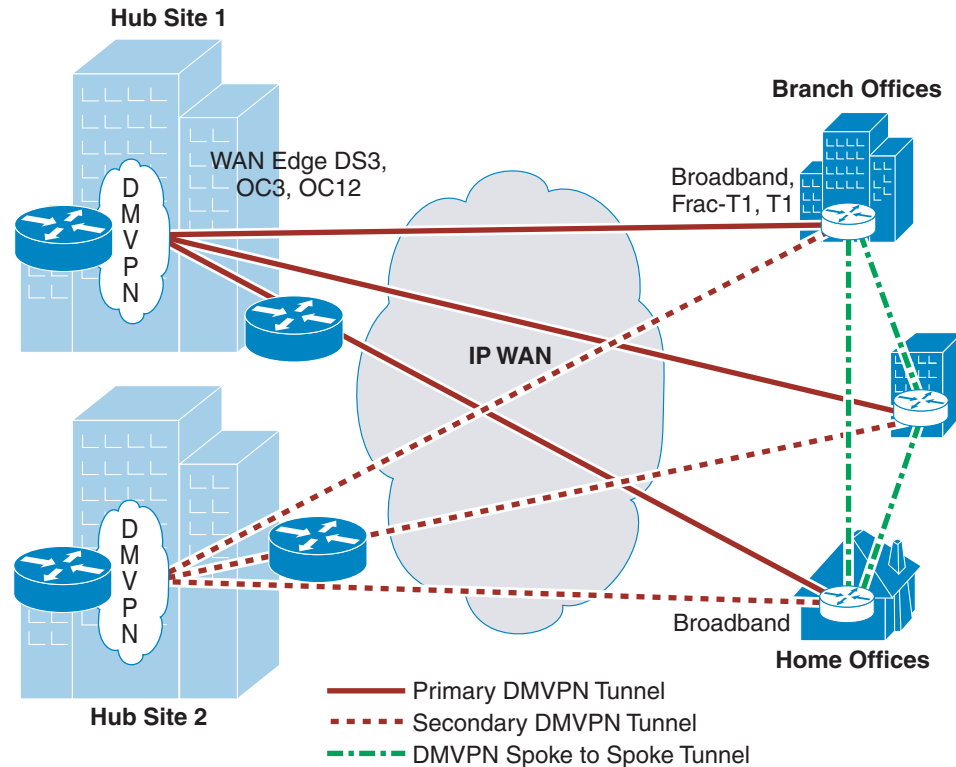


Spoke-to-Spoke Deployment Model—Headend System Architecture

A dual DMVPN cloud topology with the spoke-to-spoke deployment model supports only the Single Tier Headend Architecture. A Dual Tier Headend Architecture is not a valid option for this topology because spoke-to-spoke connections require the use of tunnel protection and IPsec profiles, which do not allow the crypto tunnel and the GRE tunnel to use different endpoint address.

Single Tier Headend Architecture

In a Single Tier Headend Architecture, the mGRE and crypto functionally co-exist on the same router CPU. [Figure 1-8](#) shows this spoke-to-spoke topology.

Figure 1-8 Single Tier Headend Architecture

In Figure 1-8, the solution is a dual DMVPN cloud topology with spoke-to-spoke deployment model. Both headends are mGRE and crypto tunnel aggregation routers servicing multiple mGRE tunnels for a prescribed number of branch office locations. In addition to terminating the VPN tunnels at the central site, headends can advertise branch routes using IP routing protocols such as EIGRP or OSPF, regardless of which DMVPN cloud path selection is chosen.

Dual DMVPN Cloud Topology—Spoke-to-Spoke Deployment Model Branch Router Considerations

Branches in a dual DMVPN cloud topology with the spoke-to-spoke deployment model provide mGRE over IPsec tunnel(s) from the branch office locations to the central site to allow the creation of branch-to-branch communication. In addition to terminating the VPN tunnels, the branch router often provides WAN access, and in some implementations may serve as a firewall.

The branch router public IP address is either a statically-defined or a dynamically-assigned IP address. Both the p2p GRE and crypto tunnels are sourced from the public IP address. This address is registered with the headend, which provides a mapping to the branch private address.

Single DMVPN Cloud Topology

In a single DMVPN cloud topology, there are two headend routers on the same DMVPN subnet. Therefore, the branch router requires an mGRE interface. Because of this mGRE interface, branch routers attempt inter-branch communications if so directed by the routing table. As a result, this model should be considered a spoke-to-spoke topology. The hub-and-spoke deployment model can be configured in a single DMVPN cloud topology with one or more headend routers. The routing protocol running on the spokes can be configured for hubs failover and/or load balancing.

A single DMVPN cloud topology with the spoke-to-spoke deployment model also contains two or more headend routers. The headend routers are configured similarly to the headend router configurations in the dual DMVPN cloud topology, but only one IP subnet is used. Traffic can be load balanced between the two headend routers, which is controlled by the routing protocol running over the DMVPN. In this topology, all branch and headend mGRE interfaces are on a single subnet, which contrasts to the dual DMVPN cloud topology where there are multiple subnets each represented by a DMVPN cloud. In this scenario, it can be slightly more complicated to configure the routing protocol to prefer one hub over the other, and it is possible that asymmetric routing between the spoke and the hubs may occur. [Figure 1-9](#) shows this deployment model.

Figure 1-9 *Single DMVPN Cloud Topology—Spoke-to-Spoke Deployment Model*



Although this is a valid topology option and Cisco does recommend this topology for spoke-to-spoke deployments, though it is not discussed in detail in this document.

Best Practices and Known Limitations

The following sections contain a summary of the best practices and limitations for the dual DMVPN cloud topology design. More detailed information is provided in [Chapter 2, “DMVPN Design and Implementation.”](#)

Best Practices Summary for Hub-and-Spoke Deployment Model

This section describes the best practices for a dual DMVPN cloud topology with the hub-and-spoke deployment, supporting IP multicast (IPmc) traffic including routing protocols.

The following are general best practices:

- Use IPsec in transport mode
- Configure Triple DES (3DES) or AES for encryption of transported data (exports of encryption algorithms to certain countries may be prohibited by law).
- Implement Dead Peer Detection (DPD) on the spokes to detect loss of communication between peers.
- Deploy hardware-acceleration of IPsec to minimize router CPU overhead, to support traffic with low latency and jitter requirements, and for the highest performance for cost.
- Keep IPsec packet fragmentation to a minimum on the customer network by setting MTU size or using Path MTU Discovery (PMTUD).
- Use Digital Certificates/Public Key Infrastructure (PKI) for scalable tunnel authentication.
- Configure a routing protocol (for example, EIGRP, BGP or OSPF) with route summarization for dynamic routing.
- Set up QoS service policies as appropriate on headend and branch router interfaces to help alleviate interface congestion issues and to attempt to keep higher priority traffic from being dropped during times of congestion.

The following are general headend best practices:

- Design the deployment to keep the headends below the critical scalability parameters for DMVPN designs:
 - Maximum number of spokes per mGRE interface
 - Maximum number of total spokes per headend

See [Chapter 3, “Scalability Considerations,”](#) for more information.

- Select Cisco VPN router products at the headend based on considerations for the following:
 - Number of tunnels to be aggregated
 - Maximum throughput in both packets per second (pps) and bits per second (bps) to be aggregated
 - Performance margin for resiliency and failover scenarios
 - Maintaining CPU utilization below design target

See [Chapter 3, “Scalability Considerations,”](#) for more information.

- Distribute branch office tunnels across a number of headend routers to balance loading and aggregation capacity of the hub(s).

The following is a Single Tier Headend Architecture best practice:

- Configure mGRE and IPsec tunnel protection on headend routers to simplify configurations and provisioning of new branches.

The following is a Dual Tier Headend Architecture best practice:

- Use dynamic crypto maps on the crypto headend to reduce the amount of IPsec configuration required.

The following are branch office best practices:

- Configure the branch with p2p GRE and IPsec tunnel protection.
- Configure two tunnels to alternate headends, using routing metrics to designate a primary and secondary path.

- Select Cisco VPN router products at the branch offices based on considerations for:
 - Maximum throughput in both pps and bps
 - Allowances for other integrated services that may be running on the router, such as for example firewall, IPS, and NAT/PAT
- See [Chapter 3, “Scalability Considerations,”](#) for more information.
- Configure **qos pre-classify** in VPN designs where both QoS and IPsec occur on the same system and QoS needs to match on parameters in the clear-text packet other than the DSCP/ToS byte. The network manager should verify correct operation.

Known Limitations Summary for Hub-and-Spoke Deployment Model

This section describes at a high level the known limitations for a dual DMVPN cloud topology with the hub-and-spoke deployment.

The following are general limitations:

- mGRE acceleration combined with the use of a GRE tunnel key is not currently supported on the Cisco Catalyst 6500/7600 router with VPN SM or VPN SPA, because neither VPN service module supports the mGRE tunnel key. These platforms can be used in designs that do not require an mGRE tunnel key. For more details, see [Chapter 2, “DMVPN Design and Implementation.”](#)
- There are significant scalability limitations for supporting IP multicast over DMVPN designs. See the *Multicast over IPsec VPN Design Guide* for more information at the following URL: http://www.cisco.com/en/US/docs/solutions/Enterprise/WAN_and_MAN/V3PNIPmc.html.
- **qos pre-classify** must be applied on the mGRE tunnel interface, because it is not currently supported by IPsec tunnel protection.

The following is a general headend limitation:

- Limited QoS can be implemented in the hub-to-branch direction on the outside interface, because it is not possible to configure a service policy at the tunnel level. This is interface level QoS, not per branch and is executed post encryption.

The following are Dual Tier Headend Architecture limitations:

- Tunnel protection is not supported.
- **qos pre-classify** is not supported in an architecture that implements two different headends for mGRE tunnels and VPN tunnels.

The following is a branch office limitation:

- Branches must always initiate the DMVPN tunnel to the headend router; the headend cannot initiate the tunnel to the branch router.

Best Practices Summary for Spoke-to-Spoke Deployment Model

This section summarizes the best practices for a dual DMVPN cloud topology with the spoke-to-spoke deployment. These best practices should be considered *in addition to* the best practices for hub-and-spoke deployments.

The following are general best practices:

- Set desired tunnel persistence timers via NHRP hold time, with consideration for IPsec SA lifetimes. For more details, see [Chapter 2, “DMVPN Design and Implementation.”](#)

- Use EIGRP routing protocols for spoke-to-spoke deployment models.

The following is a branch office best practice:

- Configure IKE Call Admission Control (IKE CAC) to limit the maximum number of spoke-to-spoke tunnels that can be accepted by a branch router, after which the tunnels go spoke-to-hub-to-spoke. For more information, see [IKE Call Admission Control, page 2-10](#).
- mGRE must be configured on the branch router.

Known Limitations Summary for Spoke-to-Spoke Deployment Model

This section describes at a high level the known limitations for a dual DMVPN cloud topology with the spoke-to-spoke deployment. These known limitations should be considered *in addition to* the known limitations for hub-and-spoke deployments.

The following are general limitations:

- ODR can only be used in spoke-to-spoke topologies if combined with VRFs (VRF-lite).
- OSPF is not recommended as a routing protocol in a spoke-to-spoke deployment model because of scaling limitations. For more information, see [Chapter 3, “Scalability Considerations.”](#)

The following is a headend limitation:

- mGRE and IPsec source and destination IP addresses must be identical for spoke-to-spoke mode to function, which is not possible with a Dual Tier Headend Architecture.

The following are branch office limitations:

- Very limited QoS can be provided between spokes. Therefore, latency-sensitive applications such as VoIP and video are considered “best effort” in spoke-to-spoke DMVPN deployments.
- Dynamic routing is not exchanged between spokes over a spoke-to-spoke tunnel. As a result, communication can be lost without knowing the tunnel is down. ISAKMP keepalives are used between the spokes to detect if the remote spoke has gone down, and if so to clear the spoke-spoke tunnel.
- Spokes behind a pNAT device must use the same outside NAT address as it used with the hubs in order to establish spoke-to-spoke tunnels.
- No IP multicast traffic can be exchanged between spokes.
- In a spoke-to-spoke topology, any traffic can bring up an IPsec tunnel to another branch in that DMVPN cloud. Because this is done at the L3 (routing) level, any IP unicast traffic can then transit over that spoke-to-spoke tunnel. This may be a security issue for some deployments because viruses, worms, or attack software may spread branch-to-branch without the headend as a check point. Other protection mechanisms such as IPS should be implemented at every branch that is spoke-to-spoke capable.
- IKE CAC has limitations as well as the maximum number of ISAKMP SA per branch platform. For more information, see [IKE Call Admission Control, page 2-10](#).

Additional detailed information on these recommendations is discussed in the chapters that follow.

