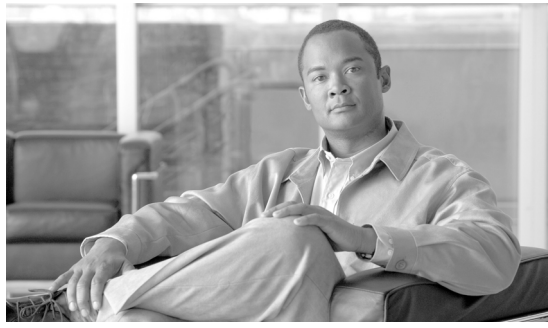




Enterprise Campus 3.0 Architecture: Overview and Framework



Note

This document is the first part of an overall systems design guide. This document will become Chapter 1 of the overall design guide when the remaining chapters are completed.

Contents

Enterprise Campus Architecture and Design Introduction	1-2
Audience	1-2
Document Objectives	1-2
Introduction	1-3
The Enterprise Campus	1-4
Campus Architecture and Design Principles	1-5
Hierarchy	1-5
Access	1-7
Distribution	1-7
Core	1-8
Mapping the Control and Data Plane to the Physical Hierarchy	1-12
Modularity	1-13
Access-Distribution Block	1-14
Services Block	1-20
Resiliency	1-22
Flexibility	1-24
Campus Services	1-25
Non-Stop High Availability	1-25
Measuring Availability	1-25



Corporate Headquarters:

Cisco Systems, Inc., 170 West Tasman Drive, San Jose, CA 95134-1706 USA

Copyright © 2008 Cisco Systems, Inc. All rights reserved.

Unified Communications Requirements	1-28
Tools and Approaches for Campus High Availability	1-30
Access and Mobility Services	1-33
Converged Wired and Wireless Campus Design	1-33
Campus Access Services	1-36
Application Optimization and Protection Services	1-38
Principles of Campus QoS Design	1-38
Network Resiliency and QoS	1-41
Virtualization Services	1-42
Campus Virtualization Mechanisms	1-43
Network Virtualization	1-44
Security Services	1-47
Infrastructure Security	1-47
Perimeter Access Control and Edge Security	1-49
Endpoint Security	1-49
Distributed Security—Defense in Depth	1-49
Operational and Management Services	1-50
Fault Management	1-51
Accounting and Performance	1-52
Configuration and Security	1-53
Evolution of the Campus Architecture	1-53

Enterprise Campus Architecture and Design Introduction

This introductory section includes the following high-level sections to present the content coverage provided in this document:

- [Audience, page 2](#)
- [Document Objectives, page 3](#)
- [Introduction, page 3](#)
- [The Enterprise Campus, page 4](#)

Audience

This document is intended for network planners, engineers, and managers for enterprise customers who are building or intend to build a large-scale campus network and require an understanding of general design requirements.

Document Objectives

This document presents an overview of the campus network architecture and includes descriptions of various design considerations, topologies, technologies, configuration design guidelines, and other considerations relevant to the design of highly available, full-service campus switching fabric. It is also intended to serve as a guide to direct readers to more specific campus design best practices and configuration examples for each of the specific design options.

Introduction

Over the last 50 years, businesses have achieved improving levels of productivity and competitive advantage through the use of communication and computing technology. The enterprise campus network has evolved over the last 20 years to become a key element in this business computing and communication infrastructure. The interrelated evolution of business and communications technology is not slowing and the environment is currently undergoing another stage of that evolution. The emerging *Human Network*, as it has been termed by the media, illustrates a significant shift in the perception of and the requirements and demands on the campus network. *The Human Network* is collaborative, interactive and focused on the real-time communications of the end-user, whoever that user may be a worker, a customer, a partner, anyone. The user experience on the network has become the critical determinant of success or failure of technology systems, whether in private or professional lives.

Web 2.0, collaborative applications, mash-ups, and the like are all reflective of a set of business and technology changes that are changing the requirements of our networking systems. An increased desire for mobility, the drive for heightened security, and the need to accurately identify and segment users, devices and networks are all being driven by the changes in the way businesses partner and work with other organizations. The list of requirements and challenges that the current generation of campus networks must address is highly diverse and includes the following:

- Global enterprise availability.
 - Unified Communications, financial, medical, and other critical systems are driving requirement for five nines (99999) availability and improved convergence times necessary for real-time interactive applications.
 - Migration towards fewer centralized data repositories increases the need for network availability for all business processes.
 - Network change windows are shrinking or being eliminated as businesses operations adjust to globalization and are operating 7x24x365.
- Collaboration and real-time communication application use is growing.
 - The user experience is becoming a top priority for business communication systems.
 - As Unified Communications deployments increase, uptime becomes even more critical.
- Continuing evolution of security threats.
 - Security threats continue to grow in number and complexity.
 - Distributed and dynamic application environments are bypassing traditional security chokepoints.
- The need to adapt to change without forklift upgrades.
 - IT purchases face longer time-in-service and must be able to adapt to adjust to future as well as present business requirements.
 - Time and resources to implement new business applications are decreasing.

- New network protocols and features are starting to appear (Microsoft is introducing IPv6 into the enterprise network).
- Expectations and requirements for anywhere; anytime access to the network are growing.
 - The need for partner and guest access is increasing as business partnerships are evolving.
 - Increased use of portable devices (laptops and PDAs) is driving the demand for full featured and secure mobility services.
 - An increasing need to support multiple device types in diverse locations.
- Next generation applications are driving higher capacity requirements.
 - Embedded rich media in documents.
 - Interactive high definition video.
- Networks are becoming more complex.
 - Do it yourself integration can delay network deployment and increase overall costs.
 - Business risk mitigation requires validated system designs.
 - Adoption of advanced technologies (voice, segmentation, security, wireless) all introduce specific requirements and changes to the base switching design and capabilities.

This document is the first part of an overall systems design guide that addresses enterprise campus architectures using the latest advanced services technologies from Cisco and is based on best-practice design principles that have been tested in an enterprise systems environment. It introduces the key architectural components and services that are necessary to deploy a highly available, secure, and service-rich campus network. It also defines a reference design framework that provides the context for each of the specific design chapters—helping the network engineer understand how specific design topics fit into the overall architecture.

The Enterprise Campus

The enterprise campus is usually understood as that portion of the computing infrastructure that provides access to network communication services and resources to end users and devices spread over a single geographic location. It might span a single floor, building or even a large group of buildings spread over an extended geographic area. Some networks will have a single campus that also acts as the core or backbone of the network and provide interconnectivity between other portions of the overall network. The campus core can often interconnect the campus access, the data center and WAN portions of the network. In the largest enterprises, there might be multiple campus sites distributed worldwide with each providing both end user access and local backbone connectivity. From a technical or network engineering perspective, the concept of a campus has also been understood to mean the high-speed Layer-2 and Layer-3 Ethernet switching portions of the network outside of the data center. While all of these definitions or concepts of what a campus network is are still valid, they no longer completely describe the set of capabilities and services that comprise the campus network today.

The campus network, as defined for the purposes of the enterprise design guides, consists of the integrated elements that comprise the set of services used by a group of users and end-station devices that all share the same high-speed switching communications fabric. These include the packet-transport services (both wired and wireless), traffic identification and control (security and application optimization), traffic monitoring and management, and overall systems management and provisioning. These basic functions are implemented in such a way as to provide and directly support the higher-level services provided by the IT organization for use by the end user community. These functions include:

- Non-Stop High Availability Services

- Access and Mobility Services
- Application Optimization and Protection Services
- Virtualization Services
- Security Services
- Operational and Management Services

In the later sections of this document, an overview of each of these services and a description of how they interoperate in a campus network is discussed. Before we look at the six services in more detail, it is useful to understand the major design criteria and design principles that shape the enterprise campus architecture. The design can be viewed from many aspects starting from the physical wiring plant, moving up through the design of the campus topology, and eventually addressing the implementation of the campus services. The order or manner in which all of these things are tied together to form a cohesive whole is determined by the use of a baseline set of design principles which, when applied correctly, provide for a solid foundation and a framework in which the upper layer services can be efficiently deployed.

Campus Architecture and Design Principles

Any successful architecture or system is based on a foundation of solid design theory and principles. Designing a campus network is no different than designing any large, complex system—such as a piece of software or even something as sophisticated as the space shuttle. The use of a guiding set of fundamental engineering principles serves to ensure that the campus design provides for the balance of availability, security, flexibility, and manageability required to meet current and future business and technological needs. The remainder of this campus design overview and related documents will leverage a common set of engineering and architectural principles: *hierarchy*, *modularity*, *resiliency*; and *flexibility*. Each of these principles is summarized in the brief sections that follow:

- [Hierarchy, page 5](#)
- [Modularity, page 13](#)
- [Resiliency, page 22](#)
- [Flexibility, page 24](#)

These are not independent principles. The successful design and implementation of an enterprise campus network requires an understanding of how each applies to the overall design and how each principle fits in the context of the others.

Hierarchy

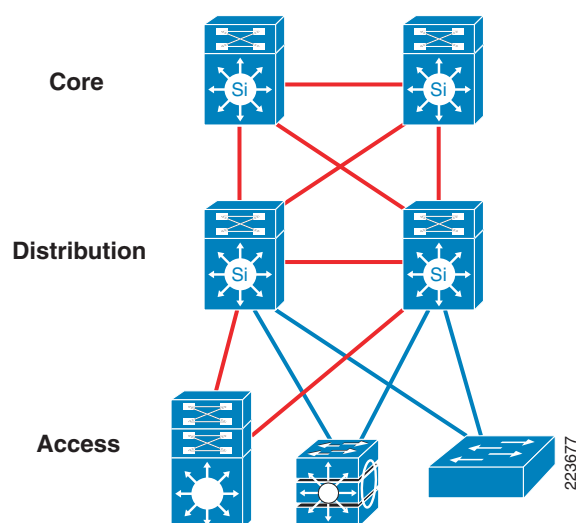
A critical factor for the successful implementation of any campus network design is to follow good structured engineering guidelines. A structured system is based on two complementary principles: *hierarchy* and *modularity*. Any large complex system must be built using a set of modularized components that can be assembled in a hierarchical and structured manner. Dividing any task or system into components provides a number of immediate benefits. Each of the components or modules can be designed with some independence from the overall design and all modules can be operated as semi-independent elements providing for overall higher system availability—as well as for simpler management and operations. Computer programmers have leveraged this principle of hierarchy and modularity for many years. In the early days of software development, programmers built *spaghetti code* systems. These early programs were highly optimized and very efficient. As the programs became larger and they had to be modified or changed, software designers very quickly learned that the lack of isolation

between various parts of the program or system meant that any small change could not be made without affecting the entire system. Early LAN-based computer networks were often developed following a similar approach. They all started as simple highly optimized connections between a small number of PCs, printers, and servers. As these LANs grew and became interconnected—forming the first generation of campus networks—the same challenges faced by the software developers became apparent to the network engineers. Problems in one area of the network very often impacted the entire network. Simple add and move changes in one area had to be carefully planned or they might affect other parts of the network. Similarly, a failure in one part of the campus quite often affected the entire campus network.

In the software development world, these sorts of system growth and complexity problems lead to the development of structured programming design using modularized or subroutine-based systems. Each individual function or software module was written in such a way that it could be changed without having to change the entire program all at once. The design of campus networks has followed the same basic engineering approach as used by software engineers. By dividing the campus system into subsystems—or building blocks—and assembling them into a clear order, we achieve a higher degree of stability, flexibility, and manageability for the individual pieces of the campus and the campus as a whole.

In looking at how structured design rules should be applied to the campus, it is useful to look at the problem from two perspectives. First, what is the overall hierarchical structure of the campus and what features and functions should be implemented at each layer of the hierarchy? Second, what are the key modules or building blocks and how do they relate to each other and work in the overall hierarchy? Starting with the basics, the campus is traditionally defined as a three-tier hierarchical model comprising the *core*, *distribution*, and *access* layers as shown in [Figure 1](#).

Figure 1 **The Layers of the Campus Hierarchy**



It is important to note that while the tiers do have specific roles in the design, there are no absolute rules for how a campus network is physically built. While it is true that many campus networks are constructed using three physical tiers of switches, this is not a strict requirement. In a smaller campus, the network might have two tiers of switches in which the core and distribution elements are combined in one physical switch, a collapsed distribution and core. On the other hand, a network may have four or more physical tiers of switches because the scale, wiring plant, and/or physical geography of the network might require that the core be extended. The important point is this—while the hierarchy of the network often defines the physical topology of the switches, they are not exactly the same thing. The key principle of the hierarchical design is that each element in the hierarchy has a specific set of functions and services that it offers and a specific role to play in each of the design.

Access

The access layer is the first tier or edge of the campus. It is the place where end devices (PCs, printers, cameras, and the like) attach to the wired portion of the campus network. It is also the place where devices that extend the network out one more level are attached—IP phones and wireless access points (APs) being the prime two key examples of devices that extend the connectivity out one more layer from the actual campus access switch. The wide variety of possible types of devices that can connect and the various services and dynamic configuration mechanisms that are necessary, make the access layer one of the most feature-rich parts of the campus network. [Table 1](#) lists examples of the types of services and capabilities that need to be defined and supported in the access layer of the network.

Table 1 *Examples of Types of Service and Capabilities*

Service Requirements	Service Features
Discovery and Configuration Services	802.1AF, CDP, LLDP, LLDP-MED
Security Services	IBNS (802.1X), (CISF): port security, DHCP snooping, DAI, IPSG
Network Identity and Access	802.1X, MAB, Web-Auth
Application Recognition Services	QoS marking, policing, queuing, deep packet inspection NBAR, etc.
Intelligent Network Control Services	PVST+, Rapid PVST+, EIGRP, OSPF, DTP, PAgP/LACP, UDLD, FlexLink, Portfast, UplinkFast, BackboneFast, LoopGuard, BPDUGuard, Port Security, RootGuard
Physical Infrastructure Services	Power over Ethernet

The access layer provides the intelligent demarcation between the network infrastructure and the computing devices that leverage that infrastructure. As such it provides a security, QoS, and policy trust boundary. It is the first layer of defense in the network security architecture and the first point of negotiation between end devices and the network infrastructure. When looking at the overall campus design, the access switch provides the majority of these access-layer services and is a key element in enabling multiple campus services.

Distribution

The distribution layer in the campus design has a unique role in that it acts as a services and control boundary between the access and the core. Both access and core are essentially dedicated special purpose layers. The access layer is dedicated to meeting the functions of end-device connectivity and the core layer is dedicated to providing non-stop connectivity across the entire campus network. The distribution layer on the other hand serves multiple purposes. It is an aggregation point for all of the access switches and acts as an integral member of the access-distribution block providing connectivity and policy services for traffic flows within the access-distribution block. It is also an element in the core of the network and participates in the core routing design. Its third role is to provide the aggregation, policy control and isolation demarcation point between the campus distribution building block and the rest of the network. Going back to the software analogy, the distribution layer defines the data input and output between the subroutine (distribution block) and the mainline (core) of the program. It defines a summarization boundary for network control plane protocols (EIGRP, OSPF, Spanning Tree) and serves as the policy boundary between the devices and data flows within the access-distribution block and the rest of the network. In providing all these functions the distribution layer participates in both the

access-distribution block and the core. As a result, the configuration choices for features in the distribution layer are often determined by the requirements of the access layer or the core layer, or by the need to act as an interface to both.

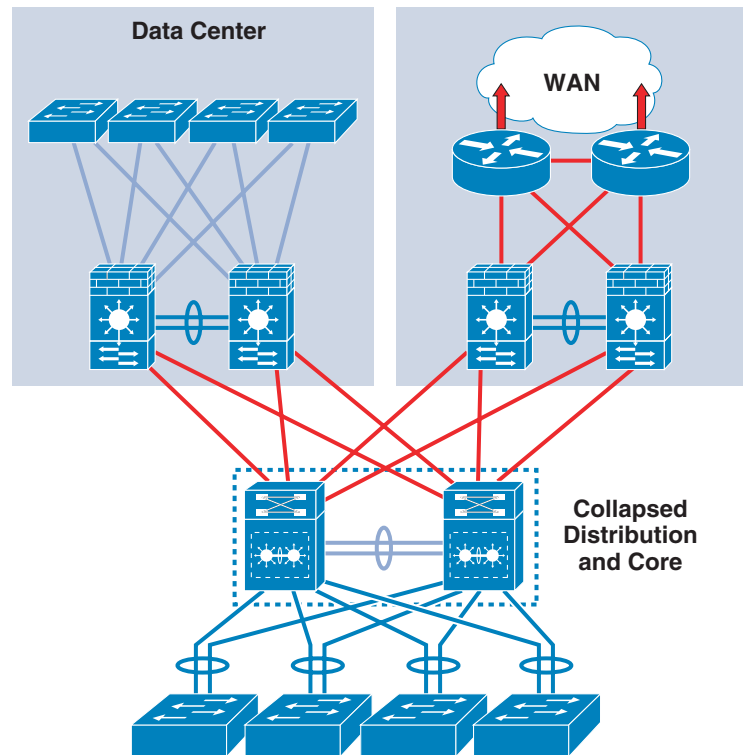
The function of the distribution layer is discussed in more detail in the description of the access-distribution block and the associated design sections.

Core

The campus core is in some ways the simplest yet most critical part of the campus. It provides a very limited set of services and is designed to be highly available and operate in an *always-on* mode. In the modern business world, the core of the network must operate as a non-stop 7x24x365 service. The key design objectives for the campus core are based on providing the appropriate level of redundancy to allow for near immediate data-flow recovery in the event of any component (switch, supervisor, line card, or fiber) failure. The network design must also permit the occasional, but necessary, hardware and software upgrade/change to be made without disrupting any network applications. The core of the network should not implement any complex policy services, nor should it have any directly attached user/server connections. The core should also have the minimal control plane configuration combined with highly available devices configured with the correct amount of physical redundancy to provide for this non-stop service capability.

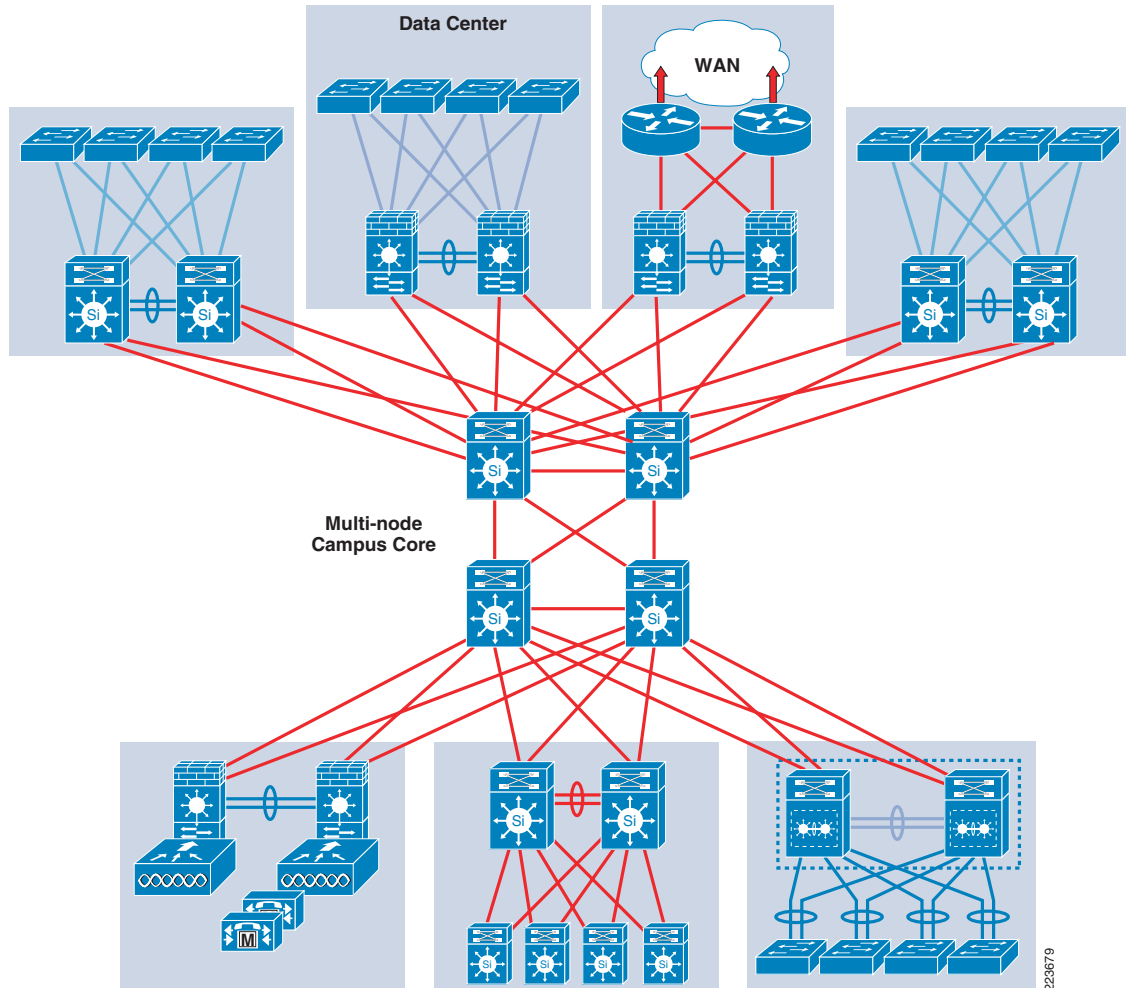
The core campus is the backbone that glues together all the elements of the campus architecture. It is that part of the network that provides for connectivity between end devices, computing, and data storage services located within the data center—and other areas and services within the network. It serves as the aggregator for all of the other campus blocks and ties together the campus with the rest of the network. One question that must be answered when developing a campus design is this: Is a distinct core layer required? In those environments where the campus is contained within a single building—or multiple adjacent buildings with the appropriate amount of fiber—it is possible to collapse the core into the two distribution switches as shown in [Figure 2](#).

Figure 2 *Collapsed Distribution and Core Campus*



It is important to consider that in any campus design even those that can physically be built with a collapsed distribution core that the primary purpose of the core is to provide fault isolation and backbone connectivity. Isolating the distribution and core into two separate modules creates a clean delineation for change control between activities affecting end stations (laptops, phones, and printers) and those that affect the data center, WAN or other parts of the network. A core layer also provides for flexibility for adapting the campus design to meet physical cabling and geographical challenges. As an example, in a multi-building campus design like that shown in [Figure 3](#), having a separate core layer allows for design solutions for cabling or other external constraints to be developed without compromising the design of the individual distribution blocks. If necessary, a separate core layer can use different transport technology, routing protocols, or switching hardware than the rest of the campus, providing for more flexible design options when needed.

Figure 3 **Multi Building Campus**

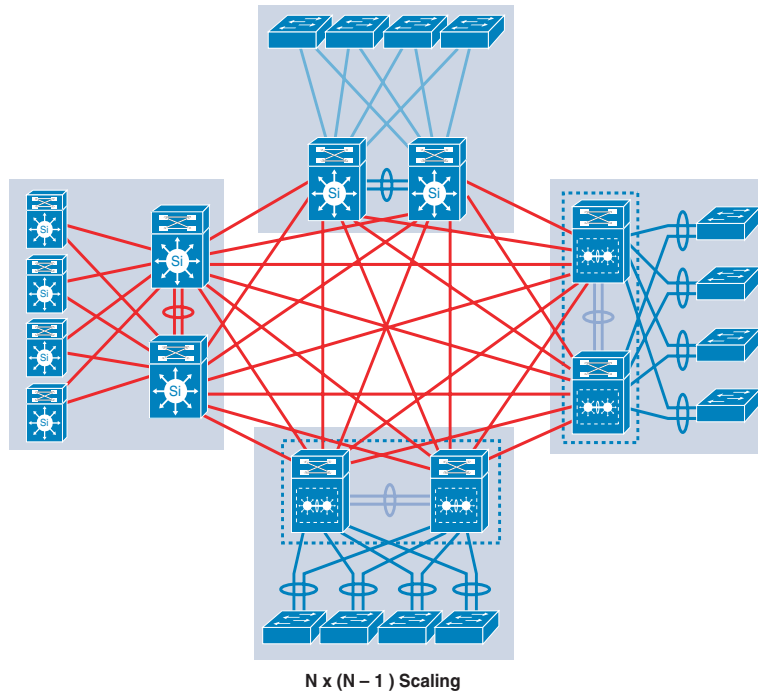


Implementing a separate core for the campus network also provides one additional specific advantage as the network grows: A separate core provides the ability to scale the size of the campus network in a structured fashion that minimizes overall complexity. It also tends to be the most cost effective solution.

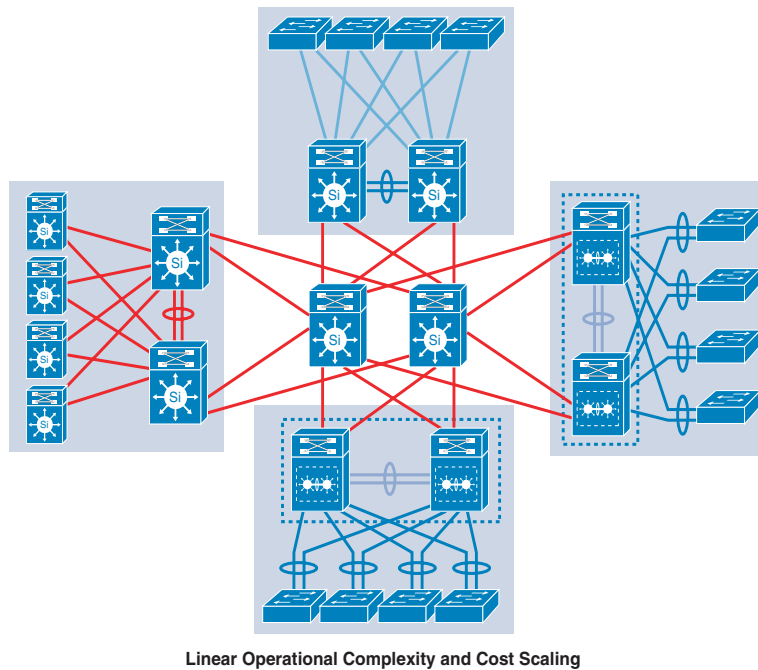
As shown in Figure 4, as the size of the network grows and the number of interconnections required to tie the campus together grow, adding a core layer significantly reduces the overall design complexity. Note that in Figure 4, the bottom design is recommended, not the top.

Figure 4 Use of Campus Core Layer to Reduce Network Scaling Complexity

Topology WITHOUT Core



Simplified Topology WITH Core



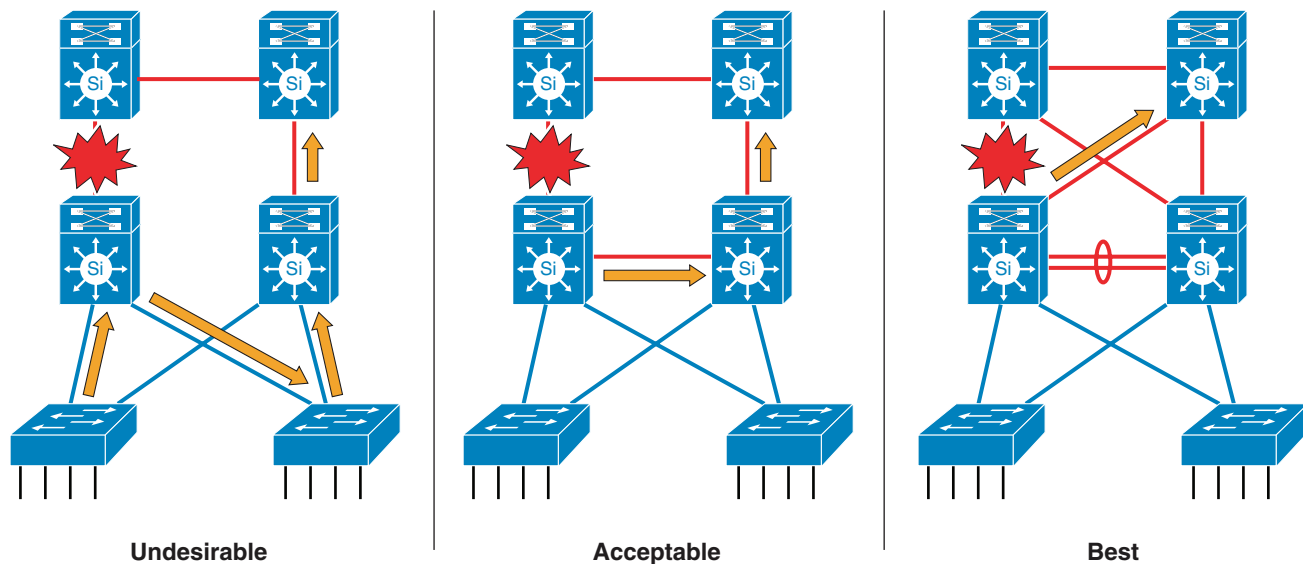
22-6680

Having a dedicated core layer allows the campus to accommodate this growth without compromising the design of the distribution blocks, the data center, and the rest of the network. This is particularly important as the size of the campus grows either in number of distribution blocks, geographical area or complexity. In a larger, more complex campus, the core provides the capacity and scaling capability for the campus as a whole.

The question of when a separate physical core is necessary depends on multiple factors. The ability of a distinct core to allow the campus to solve physical design challenges is important. However, it should be remembered that a key purpose of having a distinct campus core is to provide scalability and to minimize the risk from (and simplify) moves, adds, and changes in the campus. In general, a network that requires routine configuration changes to the core devices does not yet have the appropriate degree of design modularization. As the network increases in size or complexity and changes begin to affect the core devices, it often points out design reasons for physically separating the core and distribution functions into different physical devices.

Mapping the Control and Data Plane to the Physical Hierarchy

Implementing hierarchy in the campus network is not just a matter of physical design. In order to achieve the desired level of fault and change isolation, the logical control plane design and the data flow design must also follow hierarchical design principles. Most importantly, mapping all three elements—physical connectivity, logical control plane, and data flows—together in the same hierarchical model is necessary to produce an optimal network implementation. From a physical perspective, the distribution layer provides the boundary between the access-distribution block and the core of the network. It provides the physical demarcation between the core infrastructure and the access-distribution blocks. It should also be the demarcation and summarization point between the core control plane and the access-distribution block control plane. Having a summarized view of the connectivity and control plane within the access-distribution block allows the core and the remainder of the network to be managed and changed without constantly considering the specific internal details of the access-distribution block. The third aspect of the hierarchical design—how data traffic flows through the campus—is configured in the network, but is a desirable property or goal of the design. As shown in [Figure 5](#), the same link failure in three different switch configurations can result in three different traffic recovery paths ranging from the best case—where traffic flowing upstream recovers to another upstream path—to the worst case, in which traffic must flow back down to a lower layer of the hierarchy in order to restore network connectivity.

Figure 5 Traffic Recovery in a Hierarchical Design

223681

One of the advantages of the hierarchical design is that we can achieve a degree of specialization in each of the layers, but this specialization assumes certain network behavior. One of the assumptions or requirements that allows this specialization is that traffic is always going to flow in the same upstream or downstream hierarchical fashion (access to distribution to core). When we know that the alternative path for any traffic flow will follow the same hierarchical pattern as the original path, we can avoid making certain design decisions—such as ensuring the access layer can support extra traffic loads. Similarly, knowing that traffic always flows from the access layer through a distribution layer and then to the core, it is easier to implement consistent policy mechanisms in each layer. It reduces design complications when there is no need to consider the possibility of traffic flowing around or through a policy layer twice. Designing the hierarchy of the network to support consistent data flow behavior also has the effect of improving the network convergence time in the event of a failure. Equal-cost multi-path (ECMP) designs and other fully redundant configurations ensure these hierarchical data flows also provide for fast and deterministic convergence times over non fully meshed designs, as shown in the *Best* case in [Figure 5](#).

Modularity

The second of the two principles of structured design is *modularity*. The modules of the system are the building blocks that are assembled into the larger campus. The advantage of the modular approach is largely due to the isolation that it can provide. Failures that occur within a module can be isolated from the remainder of the network, providing for both simpler problem detection and higher overall system availability. Network changes, upgrades, or the introduction of new services can be made in a controlled and staged fashion, allowing greater flexibility in the maintenance and operation of the campus network. When a specific module no longer has sufficient capacity or is missing a new function or service, it can be updated or replaced by another module that has the same structural role in the overall hierarchical design. The campus network architecture is based on the use of two basic blocks or modules that are connected together via the core of the network:

- Access-distribution block
- Services block

The following sections introduce the underlying campus building blocks. For detailed design guidance, see each of the appropriate design document that addresses each specific module.

Access-Distribution Block

The access-distribution block (also referred to as the distribution block) is probably the most familiar element of the campus architecture. It is the fundamental component of a campus design. Properly designing the distribution block goes a long way to ensuring the success and stability of the overall architecture. The access-distribution block consists of two of the three hierarchical tiers within the multi-layer campus architecture: the access and distribution layers. While each of these layers has specific service and feature requirements, it is the network topology control plane design choices—such as routing and spanning tree protocols—that are central to determining how the distribution block glues together and fits within the overall architecture. There are currently three basic design choices for configuring the access-distribution block and the associated control plane:

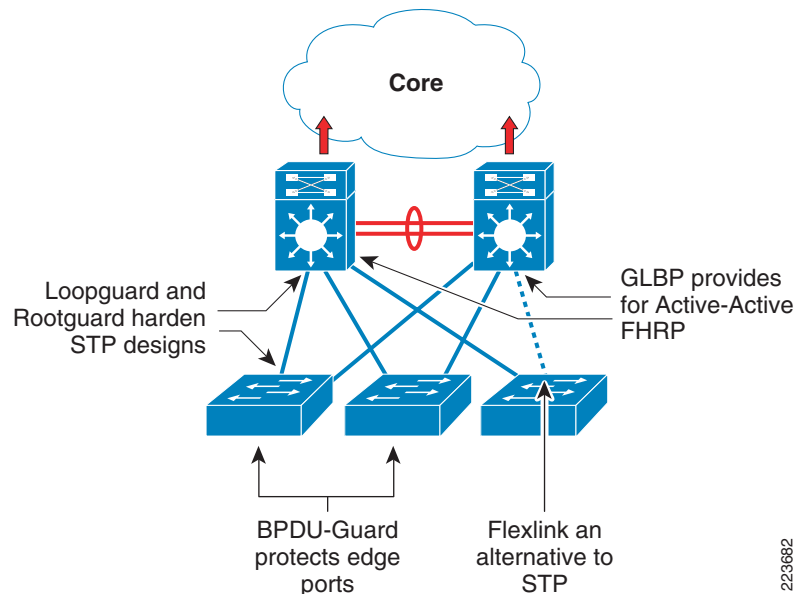
- Multi-tier
- Routed access
- Virtual switch

While all three of these designs use the same basic physical topology and cabling plant there are differences in where the Layer-2 and Layer-3 boundaries exist, how the network topology redundancy is implemented, and how load-balancing works—along with a number of other key differences between each of the design options. While a complete configuration description of each access-distribution block model can found within the detailed design documents, the following provides a short description of each design option.

Multi-Tier Access-Distribution Block

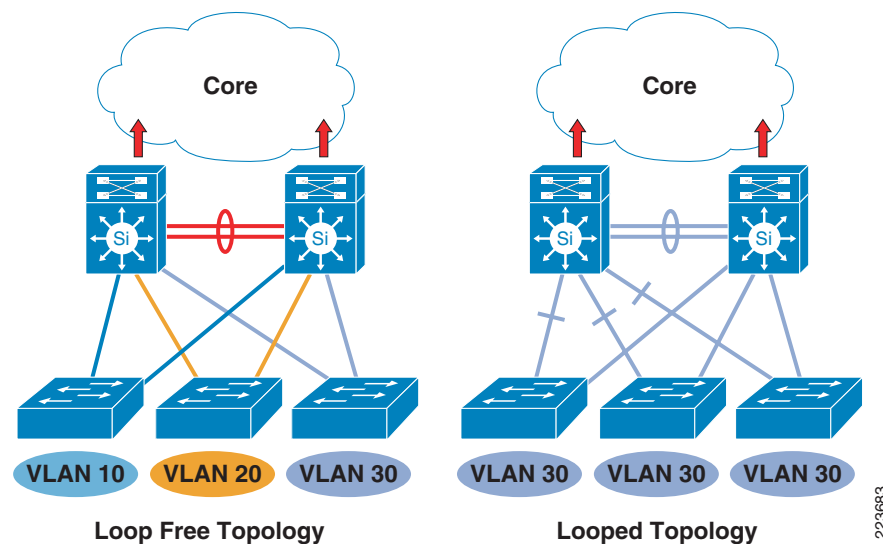
The multi-tier access-distribution model illustrated in [Figure 6](#) is the traditional campus access-distribution block design. All of the access switches are configured to run in Layer-2 forwarding mode and the distribution switches are configured to run both Layer-2 and Layer-3 forwarding. VLAN-based trunks are used to extend the subnets from the distribution switches down to the access layer. A default gateway protocol—such as HSRP or GLBP—is run on the distribution layer switches along with a routing protocol to provide upstream routing to the core of the campus. One version of spanning tree and the use of the spanning tree hardening features (such as Loopguard, Rootguard, and BPDUGuard) are configured on the access ports and switch-to-switch links as appropriate.

Figure 6 Multi-Tier Campus Access Distribution Block



The multi-tier design has two basic variations, as shown in Figure 7, that primarily differ only in the manner in which VLANs are defined. In the looped design, one-to-many VLANs are configured to span multiple access switches. As a result, each of these *spanned* VLANs has a spanning tree or Layer-2 looped topology. The other alternative—the *V* or *loop-free* design—follows the current best practice guidance for the multi-tier design and defines unique VLANs for each access switch. The removal of loops in the topology provides a number of benefits—including per device uplink load balancing with the use of GLBP, a reduced dependence on spanning tree to provide for network recovery, reduction in the risk of broadcast storms, and the ability to avoid unicast flooding (and similar design challenges associated with non-symmetrical Layer-2 and Layer-3 forwarding topologies).

Figure 7 Two Major Variations of the Multi-Tier Distribution Block

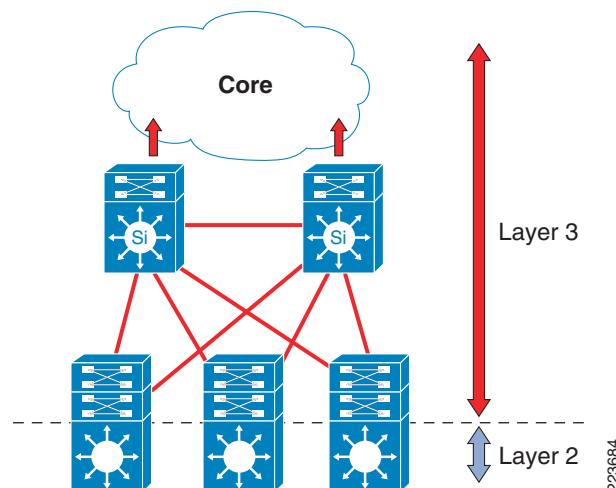


The detailed design guidance for the routed access distribution block design can be found in the campus section of the CCO SRND site <http://www.cisco.com/go/srnd>.

Routed Access Distribution Block

As alternative configuration to the traditional multi-tier distribution block model is one in which the access switch acts as a full Layer-3 routing node (provides both Layer-2 and Layer-3 switching) and the access to distribution Layer-2 uplink trunks are replaced with Layer-3 point-to-point routed links. This alternative configuration, in which the Layer-2/3 demarcation is moved from the distribution switch to the access switch appears to be a major change to the design, but is actually simply an extension of the best practice multi-tier design. See Figure 8.

Figure 8 *Routed Access Distribution Block Design*



In the best practice multi-tier and routed access design, each access switch is configured with unique voice, data, and any other required VLANs. In the routed access design, the default gateway and root bridge for these VLANs is simply moved from the distribution switch to the access switch. Addressing for all end stations and for the default gateway remains the same. VLAN and specific port configuration remains unchanged on the access switch. Router interface configuration, access lists, **ip helper** and any other configurations for each VLAN remain identical. However, these are now configured on the VLAN Switched Virtual Interface (SVI) defined on the access switch, instead of on the distribution switches. There are notable configuration changes associated with the move of the Layer-3 interface down to the access switch. It is no longer necessary to configure an HSRP or GLBP virtual gateway address, as the router interfaces for all the VLANs are now local. Similarly, with a single multicast router for each VLAN it is unnecessary to tune PIM query intervals or to ensure the designated router is synchronized with the active HSRP gateway.

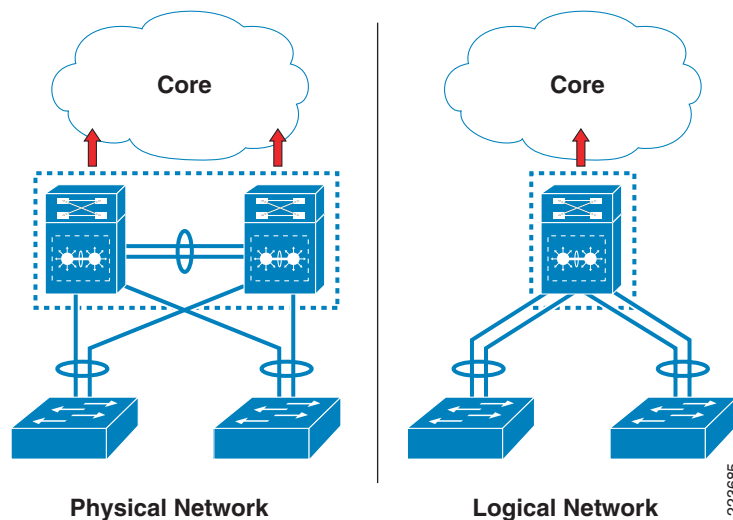
The routed access distribution block design has a number of advantages over the multi-tier design with its use of Layer-2 access to distribution uplinks. It offers common end-to-end trouble shooting tools (such as ping and traceroute), it uses a single control protocol (either EIGRP or OSPF), and removes the need for features such as HSRP. While it is the appropriate design for many environments, it is not suitable for all environments, because it requires that no VLAN span multiple access switches. The detailed design guidance for the routed access distribution block design can be found in the campus section of the CCO SRND site, <http://www.cisco.com/go/srnd>.

Virtual Switch

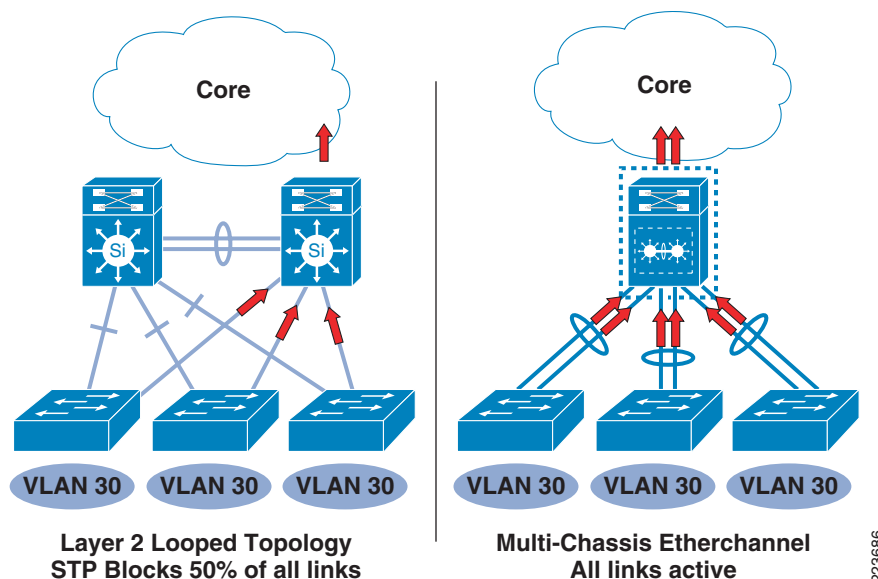
The Virtual Switching System (VSS) distribution block design is radical change from either the routed access or multi-tier designs. The introduction of the Cisco Catalyst 6500 VSS and Stackwise/Stackwise-Plus in the Cisco Catalyst 3750/3750E provides the opportunity to make a significant change to the way switch and link redundancy can be implemented. In the past, multiple

access switches were connected to two redundant distribution switches and the configuration of the network control protocols (such as HSRP, 802.1D spanning tree, and EIGRP) determined the way in which the switches forwarded traffic over each of the uplinks and the network recovered in the event of a switch or link failure. With the introduction of the virtual switch concept, the distribution switch pair can now be configured to run as a single logical switch as shown in [Figure 9](#). By converting the redundant physical distribution switches into a single logical switch, a significant change is made to the topology of the network. Rather than an access switch configured with two uplinks to two distribution switches—and needing a control protocol to determine which of the uplinks to use—now the access switch has a single multi-chassis Etherchannel (MEC) upstream link connected to a single distribution switch.

Figure 9 Virtual Switch Physical and Logical



The change from two independent uplinks to a single multi-chassis Etherchannel uplink has a number of advantages. See [Figure 10](#). Load balancing of traffic and recovery from uplink failure now leverage Etherchannel capabilities. Traffic is load-balanced per flow, rather than per client or per subnet. In the event that one of the uplinks fails, the Etherchannel automatically redistributes all traffic to the remaining links in the uplink bundle rather than waiting for spanning tree, HSRP, or other protocol to converge. The ability to remove physical Layer-2 loops from the topology—and to no longer be dependent on spanning tree to provide for topology maintenance and link redundancy—results in a distribution block design that allows for subnets and VLANs to be spanned across multiple access switches (without the traditional challenges and limitations of a spanning tree-based Layer-2 design).

Figure 10 *Virtual Switch vs. Spanning Tree Topology*

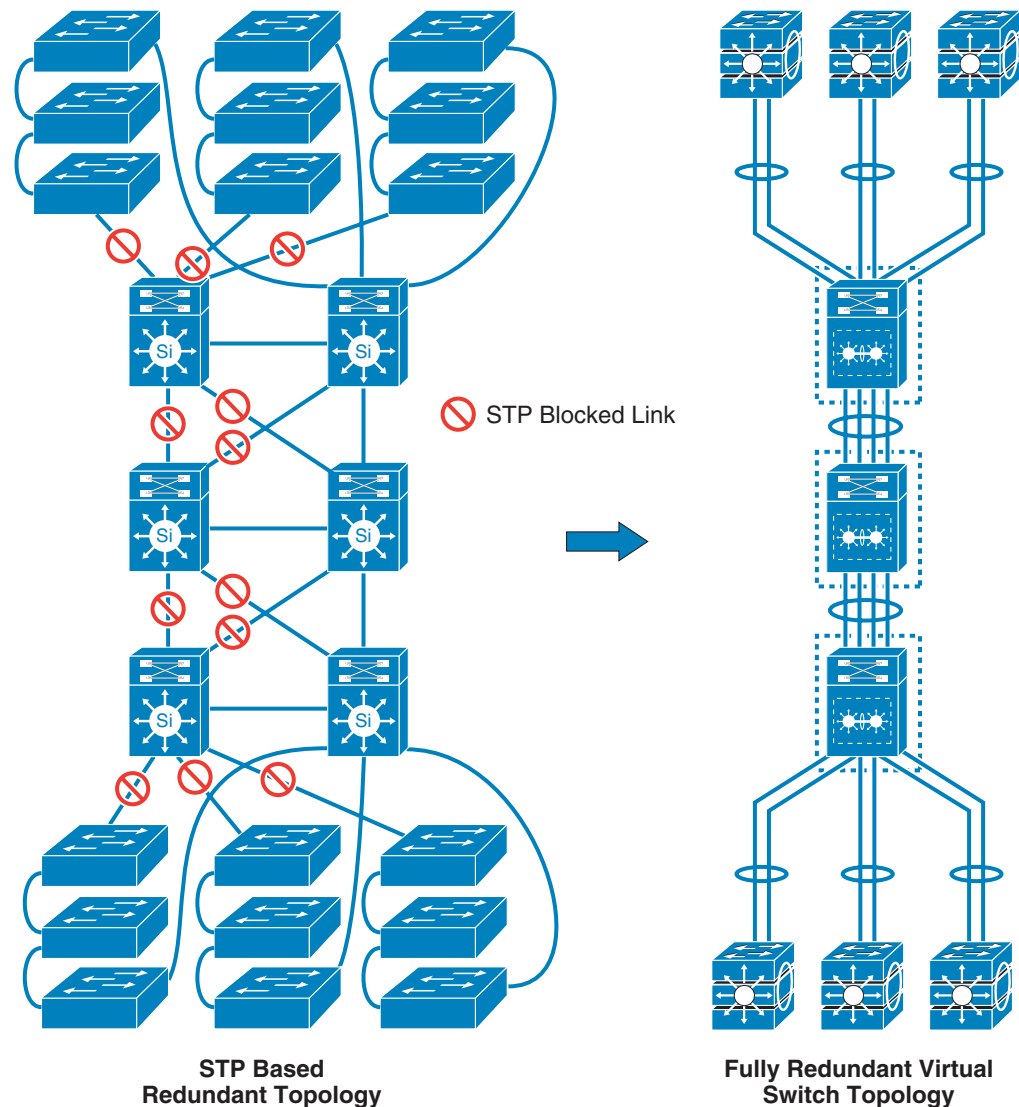
The ability to remove physical loops from the topology, and no longer be dependent on spanning tree, is one of the significant advantages of the virtual switch design. However, it is not the only difference. The virtual switch design allows for a number of fundamental changes to be made to the configuration and operation of the distribution block. By simplifying the network topology to use a single virtual distribution switch, many other aspects of the network design are either greatly simplified or, in some cases, no longer necessary. Features like HSRP or GLBP are no longer necessary because both switches act as one logical default gateway. Configuration for both per-subnet or VLAN features such as access lists, ip-helper, and others must be made only once, not replicated and kept in sync between two separate switches. Similarly, any switch configuration must be done only once and is synchronized across the redundant supervisors.

**Note**

While the virtual switch design does remove the dependency on spanning tree for active topology maintenance, spanning tree should not be turned off. Spanning tree should remain configured as a backup resiliency mechanism.

The virtual switch is not limited to the campus distribution. A virtual switch can be used in any location in the campus design where it is desirable to replace the current control plane and hardware redundancy with the simplified topology offered by the use of a virtual switch. The virtual switch simplifies the network topology by reducing the number of devices as seen by the spanning tree or routing protocol. Where two or more nodes existed with multiple independent links connecting the topology, a virtual switch can replace portions of the network with a single logical node with fewer links. [Figure 11](#) illustrates an extreme case in which an end-to-end, Layer-2 topology is being migrated from a fully redundant spanning tree-based topology to an end-to-end virtual switch-based network. Here, the topology is both drastically simplified and now all links are actively forwarding with no spanning tree loops.

Figure 11 Use of the Virtual Switch Design in an End-to-End Layer-2 Topology



While the use of a virtual switch to simplify the campus topology can help address many design challenges, the overall design must follow the hierarchical design principles. The appropriate use of Layer-2 and Layer-3 summarization, security, and QoS boundaries all apply to a virtual switch environment. Most campus environments will gain the greatest advantages of a virtual switch in the distribution layer. For details on the design of the virtual switching distribution block see the upcoming virtual switch distribution block design, <http://www.cisco.com/go/srnd>.

Distribution Block Design Comparison

While each of the three access-distribution block designs provides a viable approach, there are advantages to the virtual switch and routed access designs over the traditional multi-tier approach. Simpler overall network configuration and operation, per flow upstream and downstream load balancing, and faster convergence are some of the differences between these newer design options and the traditional multi-tier approach. The selection of a specific design option for a given campus network is

an important decision in the planning of a campus design. [Table 2](#) provides an overview comparison of the three design options. Prior to making a final design decision, review detailed design descriptions provided by Cisco to ensure that all of the factors pertinent to your environment are considered.

Table 2 **Comparison of Distribution Block Design Models**

	Multi-Tier Access	Routed Access	Virtual Switch
Access Distribution Control Plane Protocols	Spanning Tree (PVST+, Rapid-PVST+ or MST)	EIGRP or OSPF	PAgP, LACP
Spanning Tree	STP Required for network redundancy and to prevent L2 loops	No ¹	No ²
Network Recovery Mechanisms	Spanning Tree and FHRP (HSRP, GLBP, VRRP)	EIGRP or OSPF	Multi-Chassis Etherchannel (MEC)
VLAN spanning wiring closets	Supported (requires L2 spanning tree loops)	No	Supported
Layer 2/3 Demarcation	Distribution	Access	Distribution ³
First Hop Redundancy Protocol	HSRP, GLBP, VRRP required	Not Required	Not Required
Access to Distribution Per Flow Load Balancing	No	Yes - ECMP	Yes - MEC
Convergence	900 msec to 50 seconds (Dependent on STP topology and FHRP tuning)	50 to 600 msec	50 to 600 msec ⁴
Change Control	Dual distribution switch design requires manual configuration synchronization but allows for independent code upgrades and changes	Dual distribution switch design requires manual configuration synchronization but allows for independent code upgrades and changes	Single virtual switch auto-syncs the configuration between redundant hardware but does not currently allow independent code upgrades for individual member switches

1. Neither the routed access nor virtual switch designs require STP configured to maintain the network topology. It is still recommend and required to allow the use of features such as BPDU Guard on access ports.
2. Same as footnote 1.
3. With a virtual switch design, it is possible to configure a routed access layer, but this will affect the ability to span VLANs across wiring closets.
4. Initial testing indicates comparable convergence times to the routed access 50 to 600 msec. See the upcoming *Virtual Switch Design Guide* for final values.

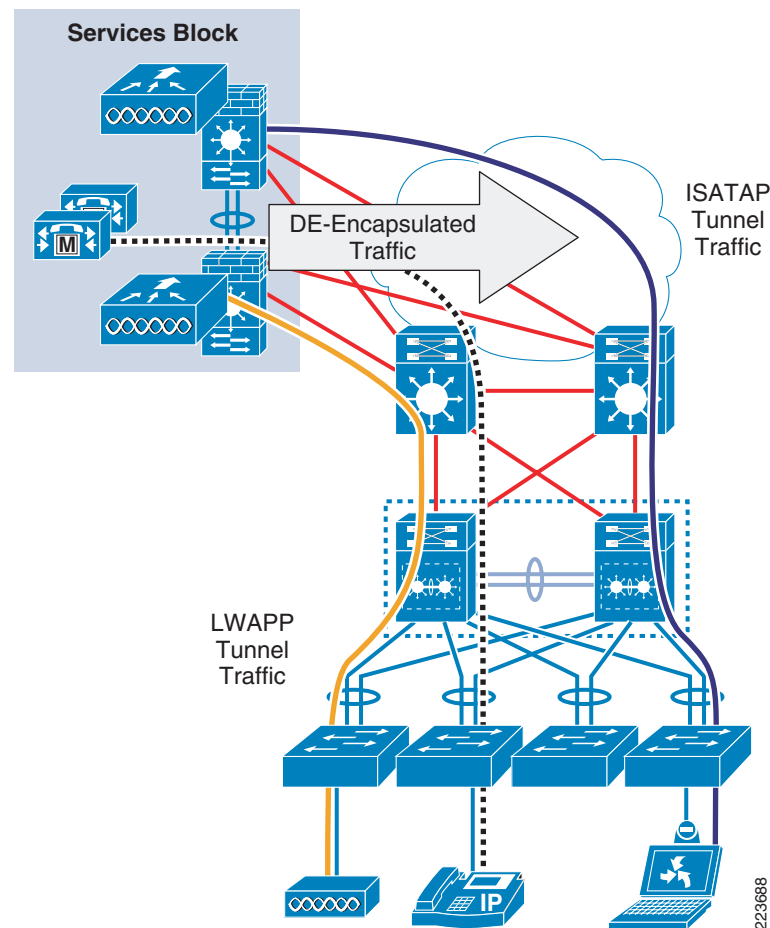
Services Block

The services block is a relatively new element to the campus design. See [Figure 12](#). As campus network planners begin to consider migration to dual stack IPv4/IPv6 environments, migrate to controller-based WLAN environments, and continue to integrate more sophisticated Unified Communications services, a number of real challenges lay ahead. It will be essential to integrate these services into the campus smoothly—while providing for the appropriate degree of operational change management and fault isolation and continuing to maintain a flexible and scalable design. As a example, IPv6 services can be deployed via an interim ISATAP overlay that allows IPv6 devices to tunnel over portions of the campus that are not yet native IPv6 enabled. Such an interim approach allows for a faster introduction of new services without requiring a network-wide, hot cutover.

Examples of functions recommended to be located in a services block include:

- Centralized LWAPP wireless controllers
- IPv6 ISATAP tunnel termination
- Local Internet edge
- Unified Communications services (Cisco Unified Communications Manager, gateways, MTP, and the like)
- Policy gateways

Figure 12 *Campus Services Block*

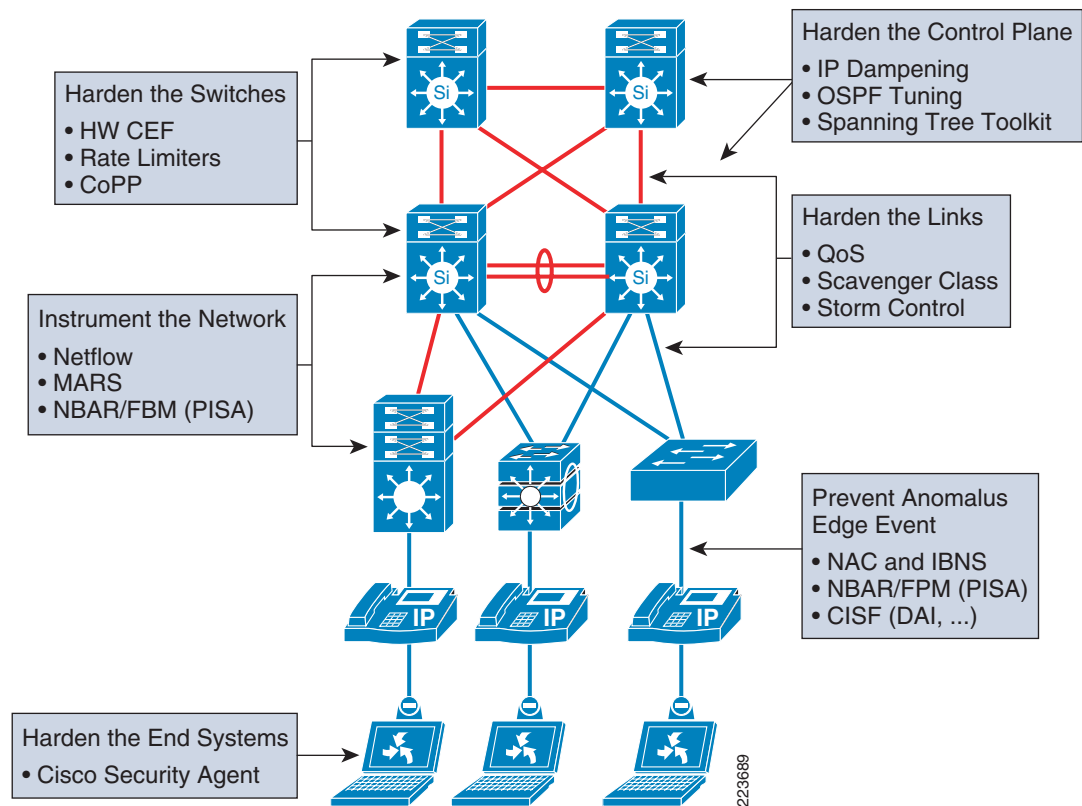


The services block is not necessarily a single entity. There might be multiple services blocks depending on the scale of the network, the level of geographic redundancy required, and other operational and physical factors. The services block serves a central purpose in the campus design; it isolates or separates specific functions into dedicated services switches allowing for cleaner operational processes and configuration management.

Resiliency

While the principles of structured design and the use of modularity and hierarchy are integral to the design of campus networks they are *not sufficient* to create a sustainable and scalable network infrastructure. Consider the software development analogy. In the software world, it is no longer sufficient for programs to merely generate the correct output given the correct input. In the same way, it is not enough that a campus network be seen as being complete solely because it correctly passes data from one point to another. As shown by the numerous security vulnerabilities exposed in software operating systems and programs in recent years, software designers are learning that to be *correct* is no longer enough. Systems must also be designed to resist failure under unusual or abnormal conditions. One of the simplest ways to break any system is to push the boundary conditions—to find the edges of the system design and look for vulnerabilities. If you are trying to break a piece of software that accepts a range of input of values from one to ten, you try giving it inputs of ten thousand, ten million, and so on to determine when and how it will break. If you are trying to break a network, follow a similar approach. Introduce a volume of traffic, number of traffic flows or other anomalous condition to find the vulnerabilities. Software engineers have become well aware of the problem and have adopted various approaches to solving it, including the use of bounds checking, assert checks, and increased modularization. Network engineers faced with a similar fundamental design challenge must also adapt network design strategies to produce a more resilient architecture.

What does it mean to create a resilient design in the context of the campus network? A basic feature of resiliency is the ability for the system to remain available for use under both normal and abnormal conditions. Normal conditions include such events as change windows and normal or expected traffic flows and traffic patterns. Abnormal conditions include hardware or software failures, extreme traffic loads, unusual traffic patterns, denial-of-service (DoS) events whether intentional or unintentional, and any other unplanned event. As illustrated in [Figure 13](#), there are a number of approaches to providing resiliency including hardening the individual components, switches, and links in the network, adding throttle or rate limiting capabilities to software and hardware functions, providing explicit controls on the behavior of edge devices, and the use of instrumentation and management tools to provide feedback to the network operations teams.

Figure 13 Examples of Campus Resiliency Features

Resilient design is not a feature nor is there a specific thing that you do in order to achieve it. As with hierarchy and modularity, resiliency is a basic principle that is made real through the use of many related features and design choices. The coordinated use of multiple features and the use of features to serve multiple purposes are aspects of resilient design. An example that illustrates this principle is the way in which an access port feature, such as port security, is used. Enabling port security on the access switch allows it to restrict which frames are permitted inbound from the client on an access port based on the source MAC address in the frame. When enabled, it can solve multiple problems—such as preventing certain man-in-the-middle and DoS flooding attacks, as well as mitigating against Layer-2 (spanning tree) loops involving the access ports. Implementing port security provides an explicit bounds check on the number of end devices that should be attached to an end port. Every network is designed to support a specific number of devices on an edge port. By implementing an explicit rule that enforces that expected behavior, the network design achieves a higher degree of overall resiliency by preventing all of the potential problems that could happen if thousands of MAC addresses suddenly appeared on an edge port. By engineering the network to both what you want it to do and prevent it from doing what you do not want it to do, you decrease the likelihood of some unexpected event from breaking or disrupting the network.

As the port security example illustrates, there are many cases where traditional security features and quality-of-service (QoS) features can and should be used to both address security and QoS requirements, but also to improve the availability of the campus infrastructure as a whole. The principle of resiliency extends to the configuration of the control plane protocols (such as EIGRP, Rapid-PVTS+, and UDLD) as well as the mechanisms used to provide switch or device level resiliency. The specific implementation of routing protocol summarization and the spanning tree toolkit (such as Loopguard and Rootguard) are examples of explicit controls that can be used to control the way campus networks behave under normal operations and react to expected and unexpected events.

Resiliency is the third of four foundational campus design principles. Just as the way in which we implement hierarchy and modularity are mutually interdependent, the way in which we achieve and implement resiliency is also tightly coupled to the overall design. Adding resiliency to the design might require the use of new features, but it is often just a matter of how we choose to implement our hierarchy and how we configure the basic Layer-2 and Layer-3 topologies.

Flexibility

In most enterprise business environments, campus networks are no longer new additions to the network. In general, campus networks have evolved through first and second generation build-out cycles and the expected lifecycle for campus networks have increased considerably—from three to five, and in some cases, seven years. At the same time, these networks have become larger and more complex, while the business environment and its underlying communication requirements continue to evolve. The result is that network designs must allow for an increasing degree of adaptability or flexibility. The ability to modify portions of the network, add new services, or increase capacity without going through a major *fork-lift* upgrade are key considerations to the effectiveness campus designs.

The structured hierarchical design inherently provides for a high degree of flexibility because it allows staged or gradual changes to each module in the network fairly independently of the others. Changes in core transport can be made independently of the distribution blocks. Changes in the design or capacity of the distribution layer can be implemented in a phased or incremental manner. Additionally, as a part of the overall hierarchical design, the introduction of the services block module into the architecture is specifically intended to address the need to implement services in a controlled fashion. This modularization of the overall design also applies to the selection of devices to fill each of the roles in the overall architecture. As the lifespan of a core, distribution, or access switch increases, it is necessary to consider how each will support and enable the continued evolution of functions required to support changing business requirements without whole scale hardware replacement.

There are a number of key areas where it is highly probable that networks will evolve over the next few years and existing designs should be adapted to incorporate the appropriate degree of flexibility into their designs to accommodate these potential changes. Key areas to consider include the following:

- *Control Plane Flexibility*—The ability to support and allow migration between multiple routing, spanning tree, and other control protocols.
- *Forwarding Plane Flexibility*—The ability to support the introduction and use of IPv6 as a parallel requirement along side IPv4.
- *User Group Flexibility*—The ability to virtualize the network forwarding capabilities and services within the campus fabric to support changes in administrative structure of the enterprise. This could involve acquisition, partnering, or outsourcing of business functions.
- *Traffic Management and Control Flexibility*—Unified communications, collaborative business approaches, and software models continue to evolve—along with a trend toward increased growth in peer-to-peer traffic flows. These fundamental changes require campus designs that allow the deployment the security, monitoring, and troubleshooting tools available to support these new traffic patterns.
- *Flexible Security Architecture*—The high probability of changing traffic patterns and a continual increase in security threats as new applications and communications patterns develop will require a security architecture that can adapt to these changing conditions.

The ability to make evolutionary modifications to any campus is a practical business and operational necessity. Ensuring that the overall architecture provides for the optimal degree of flexibility possible will ensure that future business and technology requirements will be easier and more cost effective to implement.

Campus Services

The overall campus architecture is more than the fundamental hierarchical design discussed in [Campus Architecture and Design Principles, page 5](#). While the hierarchical principles are fundamental to *how* to design a campus they do not address the underling questions about what a campus network does. What services should it provide to end users and devices? What are the expectations and parameters of those services? What functionality must be designed into each of the hierarchical layers? What must a campus network do in order to meet enterprise business and the technical requirements? What a campus does or needs to provide can be categorized into six groups:

- [Non-Stop High Availability, page 25](#)
- [Access and Mobility Services, page 33](#)
- [Application Optimization and Protection Services, page 38](#)
- [Virtualization Services, page 42](#)
- [Security Services, page 47](#)
- [Operational and Management Services, page 50](#)

In the following sections, each of these services or service level requirements is introduced. More detailed discussions of each subject will be available in the specific campus design chapters.

Non-Stop High Availability

In many cases, the principle service requirement from the campus network is the availability of the network. The ability for devices to connect and for applications to function is dependent on the availability of the campus. Availability is not a new requirement and historically has been the primary service requirement for most campus designs. The metrics of what availability means and the requirements for how *available* the network have changed as a result of the growth in unified communications, high-definition video, and the overall increasing dependence on the network for all business processes.

Measuring Availability

Availability is traditionally measured using a number of metrics, including the percentage of time the network is available or the *number of nines*—such as five nines—of availability. The calculation of availability is based on a function of the *mean time between failures* (MTBF) of the components in the network and the *mean time to repair* (MTTR)—or how long it takes to recover from a failure. See [Figure 14](#).

Figure 14 **Availability Calculation**

$$\text{Availability} = \frac{\text{MTBF}}{\text{MTBF} + \text{MTTR}}$$

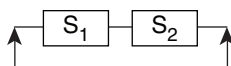
MTBF = Mean Time Between Failure
MTTR = Mean Time To Repair

223824

improving availability is achieved by either increasing the MTBF (reducing the probability of something breaking) or decreasing the MTTR (reducing the time to recover from a failure) or both. In a network with a single device this is all we need in order to consider: How reliable is the device? And how fast can we fix it if it breaks? In a network of more than one device, there are other factors that influence overall availability and our design choices.

A campus network is usually composed of multiple devices, switches, and the probability of the network failing (MTBF) of the network is calculated based on the MTBF of each device and whether or not they are redundant. In a network of three switches connected in serial, with no redundancy, the network will break if any one of the three switches breaks. The overall network MTBF is a function of how likely it is that any one of the three will fail. In a network with redundant switches, or switches in parallel, the network will only break if both of the redundant switches fail. The calculations for the system MTBF are based on the probability that one switch in a non-redundant (serial) network breaks ([Figure 15](#)), or both switches in a redundant (parallel) design break ([Figure 16](#)).

Figure 15 MTBF Calculation with Serial Switches



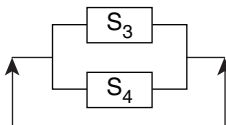
S₁, S₂ - Series Components

System is available when both components are available:

$$A_{\text{series}} = A_1 \times A_2$$

223825

Figure 16 MTBF Calculation with Parallel Switches



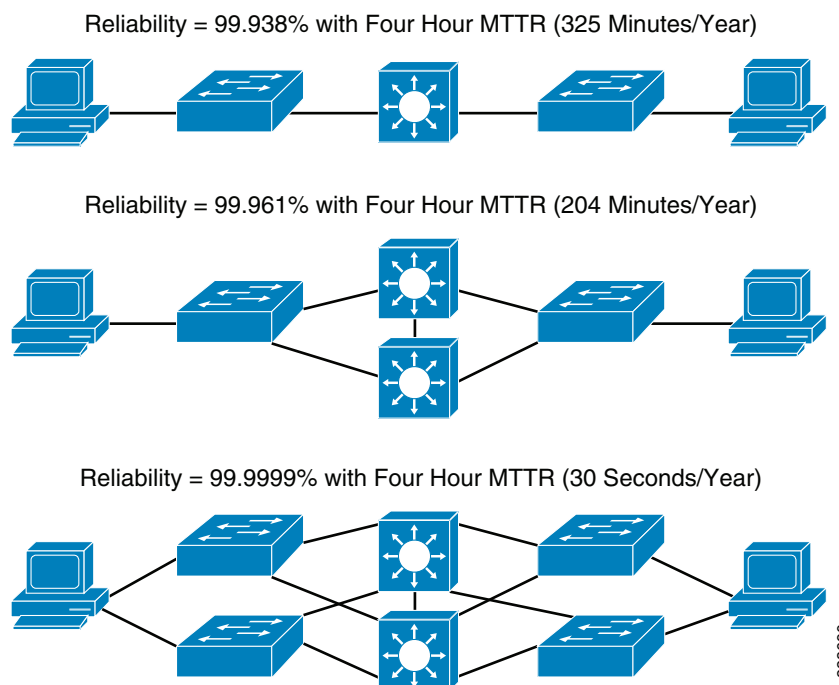
S₃, S₄ - Parallel Components

System is unavailable when both components are unavailable:

$$A_{\text{parallel}} = 1 - (1 - A_1) \times (1 - A_2)$$

223826

In addition to changing the MTBF calculations, redundancy and how redundancy is used in a design also affects the MTTR for the network. See [Figure 17](#). The time to restore service, data flows, in the network is based on the time it takes for the failed device to be replaced or for the network to recover data flows via a redundant path. The time it takes any operations team to replace a device is usually measured in hours or days rather than in minutes or seconds and the impact on the availability of the network can be significant if the appropriate degree of device redundancy is missing from the design.

Figure 17 Impact of network redundancy on overall campus reliability

The other commonly used metric for measuring availability is *defects per million* (DPM). While measuring the probability of failure of a network and establishing the service-level agreement (SLA) that a specific design is able to achieve is a useful tool, DPM takes a different approach. It measures the impact of defects on the service from the end user perspective. It is often a better metric for determining the availability of the network because it better reflects the user experience relative to event effects. DPM is calculated based on taking the total *affected user minutes* for each event, total users affected, and the duration of the event, as compared to the total number of service minutes available during the period in question. You divide the sum of service downtime minutes by total service minutes and multiply by 1,000,000. See [Figure 18](#).

Figure 18 Defects per Million Calculation

$$DPM = \frac{\sum(\#ofUsers_Affected \times Outage_Minutes)}{(\#Total_Users \times Total_Service_Minutes)} \times 10^6$$

223827

DPM is useful in that it is a measure of the observed availability and considers the impact to the end user as well as the network itself. Adding this user experience element to the question of campus availability is very important to understand and is becoming a more important part of the question of what makes a highly available or non-stop campus network. A *five nines* network, which has been considered the hallmark of excellent enterprise network design for many years, allows for up to five (5) minutes of outage or downtime per year. See [Table 3](#).

Table 3 Availability, DPM and Downtime

Availability (Percent)	DPM	Downtime/Year (24x7x365)		
99.000	10,000	3 Days	15 Hours	36 Minutes
99.500	5,000	1 Day	19 Hours	48 Minutes

Table 3 **Availability, DPM and Downtime (continued)**

Availability (Percent)	DPM	Downtime/Year (24x7x365)		
99.900	1,000		8 Hours	46 Minutes
99.950	500		4 Hours	23 Minutes
99.990	100			53 Minutes
99.999	10			5 Minutes
99.9999	1			0.5 Minutes

From a network operations perspective, achieving a maximum of five minutes of downtime over the year is a significant goal. However, as a single metric, it is not sufficient to characterize a network as meeting the availability requirements of the current and evolving business environments. DPM takes into consideration the measurement of the availability of the network from the user (or application) perspective and is a valuable tool to determine whether or not the network SLA is being met. Nonetheless, it is not a sufficient metric either. The third metric to be considered in the campus design is the *maximum outage* that any application or data stream will experience during a network failure. Network recovery time from the user (or application) perspective is the third critical design metric to consider when designing a campus network. Five minutes of outage experienced in the middle of a critical business event has a significant impact on the enterprise.

Unified Communications Requirements

Providing for a high availability in a campus design requires consideration of three aspects:

- What SLA can the design support (how many nines)?
- Is the network meeting the SLA (DPM)?
- What will the impact of any failure be on applications and user experience?

The first two are aggregated metrics of the operational integrity of a campus network and are used to determine the level of operational reliability of the network. The third consideration is a measure of business disruption—how disruptive to the business will any failure be. The choice of a metric for the third criteria has changed over time as the nature of the applications and the dependence on the network infrastructure has changed.

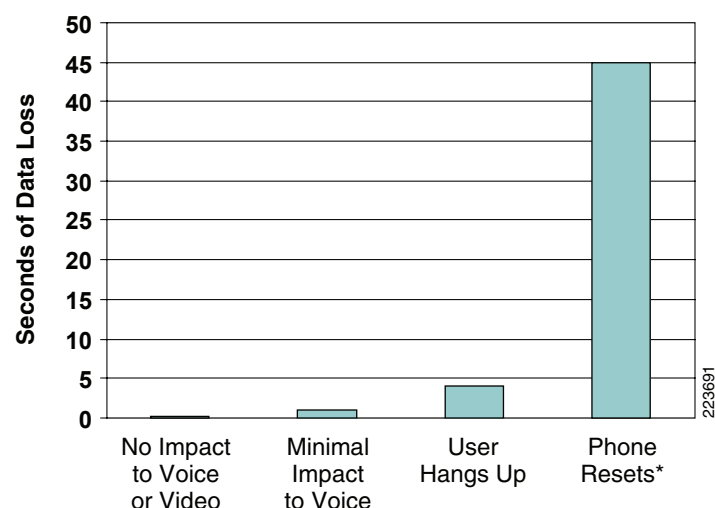
As enterprises migrate to VoIP and Unified Communications, what is considered acceptable availability must also be re-evaluated. The upper limit for acceptable network reconvergence, the MTTR, for a Unified Communications must consider several key metrics:

- How fast must the network restore data flows before the loss becomes disruptive to an interactive voice or video? When will your conversation be disrupted?
- How fast must the network converge and restore data flows before someone hangs up on an active conversation due to dead air? How long will someone listen to the phone if they do not hear anything? How long will it be before the network appears broken?
- How fast must the network converge to avoid call signalling failures, loss of dial tone, reset triggered by loss of connection to the call agent (such as Cisco Unified Communications Manager, Cisco Unified SRST, or Cisco Unified Communications Manager Express)?

These metrics contain objective and subjective elements. In addition to defining when applications will fail, they also define what is disruptive to the employees and users of the network, what events will disrupt their ability to conduct business, and what events signify a failure of the network. As network-based communications become the norm for all aspects of personal and business life, the defining of metrics describing a *working* network is increasingly important and more restrictive.

While the metrics to evaluate subjective failure assessment are by definition subjective, they do have a basis in the common patterns of human communication patterns. The amount of time that a person is willing to listen to dead air before deciding that the call (network) failed—causing the user to hang up—is variable, but tends to be in the 3-to-6 second range. The length of data or bearer path loss in an RTP stream is much stricter. While the human ear can detect loss of sound in streaming audio down to 50 msec or less, the average interval that proves disruptive to a conversation is closer to 200 msec. The ability to fill lost phonetic information in a conversation and the threshold for what period of time constitutes a pause in speech—signalling it is someone else’s turn to talk—are much longer than what the human ear can detect as lost sound. Loss of sound for periods of up to one second are recovered in normal speech pattern relatively easily, but beyond that they become disruptive to conversation and result in lost or failed communication. See [Figure 19](#).

Figure 19 Comparative Measure of MTTR on Unified Communications



*The time for a phone to reset is variable and depends on the signaling protocol, SCCP or SIP, and the state of the call, active, ringing, ...

A campus that can restore RTP media streams in less time than it takes to disrupt an active business conversation is as much a design objective in a Unified Communications-enabled enterprise as is meeting a target of five nines of availability.



Note

Voice and video are not the only applications with strict convergence requirements. Trading systems, health care, and other real-time applications might have just as strict or even more strict requirements for network recovery speed. Voice is used as a metric for the Cisco enterprise design guides because it is becoming a standard application in most enterprise networks and provides a common objective that all designs must meet as a minimum requirement.

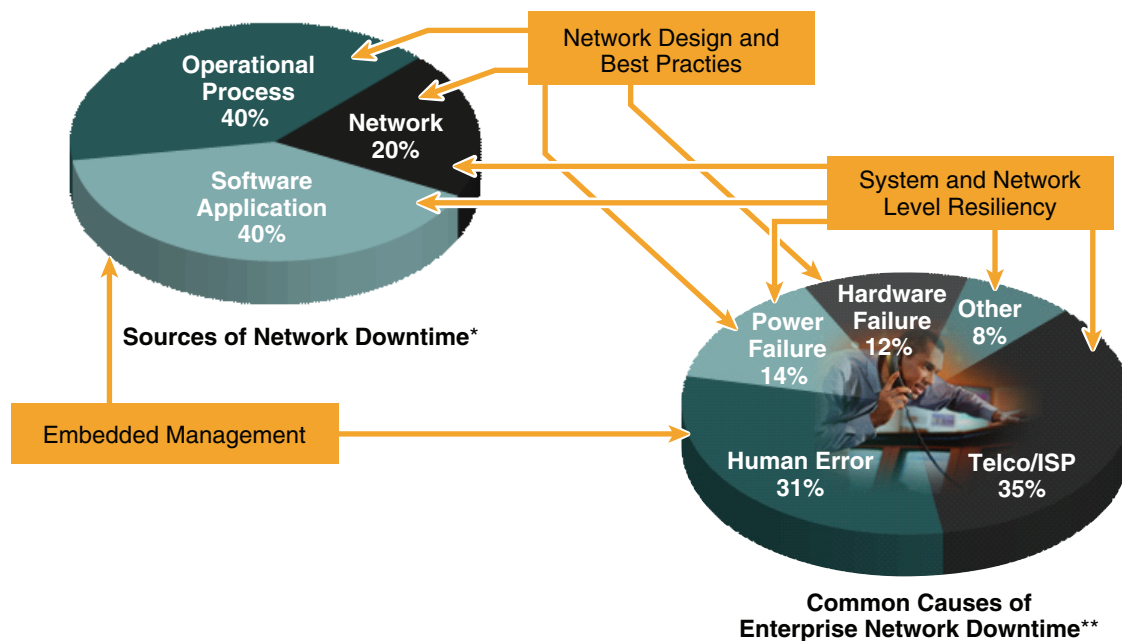
Tools and Approaches for Campus High Availability

The approach taken in the ESE campus design guide to solving both the problem of ensuring five nines of availability and providing for the recovery times required by a Unified Communications-enabled campus is based on approaching the high-availability service problem from three perspectives:

- Network resiliency
- Device resiliency
- Operational resiliency

This approach is based on an analysis of the major contributing factors of network downtime (as illustrated in Figure 20) and by using the principles of hierarchy, resiliency, and modularity—combined with the capabilities of the Cisco Catalyst switching family to define a set of design recommendations.

Figure 20 Common Causes of Network Downtime



*Source: Gartner Group

**Source: Yankee Group The Road to a Five-Nines Network 2/2004

223692

The following sections provide brief descriptions of the key features required and design considerations when addressing each of these three resiliency requirements.

Network Resiliency

Network resiliency is largely concerned with how the overall design implements topology redundancy, redundant links and devices, and how the control plane protocols (such as EIGRP, OSPF, PIM, and STP) are optimally configured to operate in that design. The use of physical redundancy is a critical part of ensuring the availability of the overall network. In the event of a component failure, having a redundant component means the overall network can continue to operate. The control plane capabilities of the campus provide the ability to manage the way in which the physical redundancy is leveraged, the network load balances traffic, the network converges, and the network is operated. The detailed recommendations for how to optimally configure the various control plane protocols are covered in the specific campus design guide, but the following basic principles can be applied in all situations:

- Wherever possible, leverage the ability of the switch hardware to provide the primary detection and recovery mechanism for network failures (for example, use Multi-Chassis Etherchannel, Equal Cost Multi-Path recovery for failure recovery). This ensures both a faster and a more deterministic failure recovery.
- Implement a *defense-in-depth* approach to failure detection and recovery mechanisms. An example of this is configuring the UniDirectional Link Detection (UDLD) protocol which uses a Layer-2 keep-alive to test that the switch-to-switch links are connected and operating correctly and acts as a backup to the native Layer-1 unidirectional link detection capabilities provided by 802.3z and 802.3ae standards.
- Ensure that the design is self-stabilizing. Utilize a combination of control plane modularization (such as route summarization) and software throttling (such as IP interface dampening) to ensure that any failures are isolated in their impact and that control plane prevents any flooding or thrashing conditions from arising.

These principles are intended to be a complementary part of the overall structured modular design approach to the campus architecture and primarily serve to re-enforce good resilient design practices.

Device Resiliency

While a redundant network topology, featuring redundant links and switches, can help address many overall campus availability challenges, providing redundancy alone does not comprise a complete solution. Every campus design will have single points of failure and the overall availability of the network might be dependent on the availability of a single device. A prime example of this is the access layer. Every access switch represents a single point of failure for all of the attached devices. Ensuring the availability of the network services is often dependent on the resiliency of the individual devices.

Device resiliency, as with network resiliency, is achieved through a combination of the appropriate level of physical redundancy, device hardening, and supporting software features. Studies indicate that most common failures in campus networks are associated with Layer-1 failures—from components such as power, fans, and fiber links. The use of diverse fiber paths with redundant links and line cards combined with fully redundant power supplies and power circuits, are the most critical aspects of device resiliency. The use of redundant power supplies becomes even more critical in access switches with the introduction of Power over Ethernet (PoE) devices such as IP phones. Multiple devices are now dependent on the availability of the access switch and its ability to maintain the necessary level of power for all of the attached end devices. After physical failures, the most common cause of device outage is often related to the failure of supervisor hardware or software. The network outages due to the loss or reset of a device due to supervisor failure can be addressed through the use of supervisor redundancy. Cisco Catalyst switches provides two mechanisms to achieve this additional level of redundancy:

- Stateful switchover and non-stop forwarding (NSF/SSO) on the Cisco Catalyst 4500 and Cisco Catalyst 6500
- Stackwise and Stackwise-Plus on the Cisco Catalyst 3750 and Cisco Catalyst 3750E

Both of these mechanisms provide for a hot active backup for the switching fabric and control plane—ensuring that both data forwarding and network control plane (featuring protocols such as EIGRP, OSPF, and STP) seamlessly recover (sub-second traffic loss) during any form of software or supervisor hardware crash.



Note

For additional information on improving the device resiliency in your campus design see the Campus Redundant Supervisor Design chapter.

In addition to ensuring that each switch in the campus has the necessary level of physical hardware and software redundancy, it is also important to provide the appropriate protection for the switches control plane. The multi-gigabit speeds of modern switching networks can overwhelm the capacity of any CPU. While most traffic in the campus network is forwarded in the hardware and the CPU should only need to process control plane and other systems management traffic, the potential exists under certain failure conditions (or in the event of a malicious DoS attack) for the volume and type of traffic forwarded to overwhelm the CPU. In such events, unless the appropriate switch hardware architecture and controls are in place, the network as a whole can fail due to the CPU being unable to process critical control plane (e.g., EIGRP and STP) and management (such as Telnet and SSH) traffic. The campus design addresses this type of problem through three approaches:

- Limit the baseline control plane and CPU load on each switch through modular design, as well as to provide control plane isolation between modules in the event any failure does occur.
- Reduce the probability of a flooding event through the reduction in the scope of the Layer-2 topology and the use of the spanning tree toolkit features to harden the spanning tree design.
- Leverage the hardware CPU protection mechanisms and Control Plane Protection (CoPP) features of the Catalyst switches to limit and prioritize traffic forwarded to each switch CPU.

The combination of all three elements (physical redundancy to address Layer-1 physical failures, supervisor redundancy to provide for a non-stop forwarding (data) plane, and the hardening of the control plane through the combination of good design and hardware CPU protection capabilities) are the key elements in ensuring the availability of the switches themselves and optimal uptime for the campus as a whole.

Operational Resiliency

Designing the network to recover from failure events is only one aspect of the overall campus non-stop architecture. Business environments are continuing to move toward requiring true 7x24x365 availability.

It is becoming increasingly difficult to find a *change window*—or a time when the network can be shut down for maintenance with the globalization of business, the desire for *always-on* communications and the movement from mainframe-based monolithic application systems to web- and Unified Communications-based systems.

The campus—which might form or be a part of the backbone of the enterprise network—must be designed to enable standard operational processes, configuration changes, software and hardware upgrades without disrupting network services.

The ability to make changes, upgrade software, and replace or upgrade hardware in a production is possible due to the implementation of network and device redundancy. By having dual active paths through redundant switches designed to converge in sub-second timeframes, it is possible to schedule an outage event on one element of the network and allow it to be upgraded and then brought back into service with minimal disruption to the network as a whole. The ability to upgrade individual devices without taking them out of service is similarly based on having internal component redundancy (such as with power supplies, and supervisors) complemented with the system software capabilities. Two primary mechanisms exist to upgrade software in place in the campus:

- Full-image In-Service Software Upgrade (ISSU) on the Cisco Catalyst 4500 leverages dual supervisors to allow for an full, in-place Cisco IOS upgrade. Moving from 12.2(37)SG1 to 12.2(40)SG, as an example. This leverages the NSF/SSO capabilities of the switch and provides for less than 200 msec of traffic loss during a full Cisco IOS upgrade.
- Sub-system ISSU on the Cisco Catalyst 6500 leverages Cisco IOS modularity and the ability it provides to replace individual Cisco IOS components (such as routing protocols) without impacting the forwarding of traffic or other components in the system.

Having the ability to operate the campus as a non-stop system is dependent on the appropriate capabilities being designed-in from the start. Network and device level redundancy, along with the necessary software control mechanisms, guarantee controlled and fast recovery of all data flows following any network failure—while concurrently providing the ability to proactively manage the non-stop infrastructure.

Access and Mobility Services

Of all the factors influencing change in the campus architecture, the growing expectation within the business community for a flexible work environment—providing anytime/anywhere network connectivity—is one of the most visible. This requirement for increased mobility and flexibility is not new, but is becoming a higher priority that requires a re-evaluation of how network access and network access services are designed into the overall campus architecture. The growth in demand for enhanced mobility—both wired and wireless—can be characterized by observing three loosely related trends:

- The growth in laptop and other portable devices as the primary business tool rather than desktop PCs.
- The growth in the number of onsite partners, contractors and other guests using the campus services. These users will most often leverage a combination of their own computing equipment—usually their corporate provided laptop—and equipment, phones, printers, and the like provided by the host enterprise.
- The growth in the number and type of devices connected to the campus network, such as VoIP phones, desktop video cameras, and security cameras.

The single thread that ties all of the requirements together is the need to cost-effectively move devices within the campus and have them associated with the correct network policies and services wherever they are connected. In order to achieve this level of access mobility, the campus network must ensure that the following access services are integrated into the overall campus architecture:

- Ability to physically attach to the network and be associated with or negotiate the correct Layer-1 and Layer-2 network services—PoE, link speed and duplex, subnet (VLAN or SSID)
- Ability to provide device identification and, where needed, perform network access authentication
- Ability for the network to apply the desired QoS policies for the specific user, device or traffic flow (such as RTP streams)
- Ability for the network to apply the desired security policies for the specific user or device
- Ability for the network and device to determine and then register the location of the attaching device
- Ability for the device to negotiate and register the correct end station parameters (such as DHCP), as well as register for any other necessary network services (such as register for Unified Communications presence and call agent services)

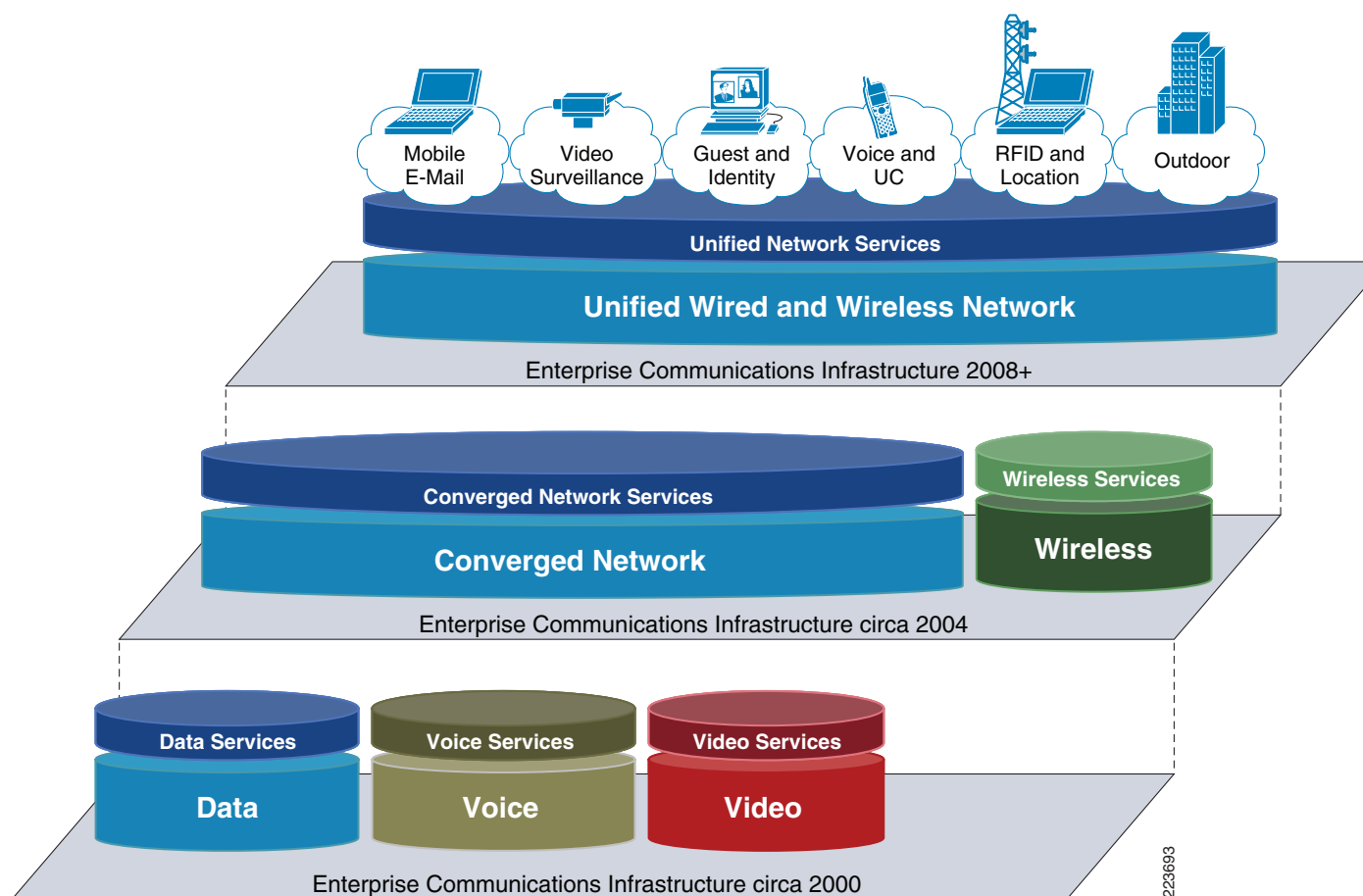
The challenge for the campus architect is determining how to implement a design that meets this wide variety of requirements, the need for various levels of mobility, the need for a cost-effective and flexible operations environment, while being able to provide the appropriate balance of security and availability expected in more traditional, fixed-configuration environments.

Converged Wired and Wireless Campus Design

One approach that is being used to address this growing need for more dynamic and flexible network access is the introduction of 802.11 wireless capabilities into the campus. While 802.11 can and does provide for easier roaming and can provide a cost effective method to enhance network access, the implementation of wireless must be integrated into an overall campus architecture in order to provide

for a consistent set of services and ease of movement for both highly mobile wireless devices and highly available wired devices. The integration of wired and wireless access methods into a common campus architecture is just the latest phase of network convergence. As illustrated in Figure 21 (moving from the bottom to the top) the enterprise network has gone through several phases of integration or convergence.

Figure 21 *Evolution of the Converged Campus Networks*



There are two key motivators that have been driving the network convergence process. The first is the ability for a converged network to reduce the operational costs of the overall enterprise by leveraging common systems and (more importantly) a common operational support teams and processes. The second, and equally important, driver to convergence is the business advantage gained when previously isolated business processes can be more tightly integrated. The convergence of the voice, video, and data networks (as an example) has enabled the development of Unified Communications systems that are allowing businesses to more efficiently leverage all the various inter-personal communication tools. This next phase of integration, combining wired and wireless into a converged campus, is motivated by the same reasons. Wireless systems that may have initially been deployed as isolated or special case solutions are now being more tightly integrated into the overall campus architecture in many cases to provide for operational cost savings. Leveraging common authentication backend systems, desktop clients, common security services, and the like—along with the use of common support processes—can result in a more efficient and effective operational environment. Just as importantly, the ability to provide

business efficiencies by being able to seamlessly move a device between wired and wireless environments and to provide for collaboration and common services between devices independent of underlying physical access connectivity type is a key requirement for this next phase of converged design.

As a part of the process of developing the overall converged wired and wireless access architecture, it is important to understand that the drive to provide enhanced mobility must be balanced with the need to support mission critical applications. Currently there are still differences in the properties and capabilities of the wired and wireless access technologies that need to be analyzed when deciding which devices should utilize wired, which should use wireless, and which need the ability to move back and forth based on changing requirements. The decision matrix used to determine when a device should be configured to use wired access versus wireless access has a number of specific factors, but it essentially distills down into a question of where a device and its application requirements sits on a spectrum of strict service level requirement versus ease-of-mobility. See [Figure 22](#).

Figure 22 **Wired vs. Wireless Decision Keys**



One of the key differences between wired and wireless environments is primarily a function of the differences between shared and dedicated media. The wired access port is a switched full duplex resource with dedicated hardware resources providing the access services (QoS, security) for each client. The wireless media is a shared resource that leverages arbitration protocols to allocate fair usage of the shared media. The result of this basic difference is that while wireless access provides for a highly flexible environment allowing seamless roaming throughout the campus it suffers the risk that the network service will degrade under extreme conditions and will not always be able to guarantee network service level requirements. Wired ports provide much more reliable guarantees for QoS (jitter, latency), packet reliability (multicast), and offer much higher capacity and fundamentally more isolation for Layer-1 and Layer-2 problems. However, a wired port is a fixed-location resource. [Table 4](#) provides a breakdown of some decision criteria that can be used to evaluate the tradeoffs between wired vs. wireless access.

Table 4 **Comparison of Wired vs. Wireless Support of Application Requirements**

	Wired	Wireless
Availability	Switched Ethernet provides for inherent layer 1 fault isolation and when complimented by capabilities in the current Catalyst switches provides for layer 2 fault isolation and DoS protection. ¹	Modern 5Ghz WLAN systems with centralized radio management provide multiple layers of protection against radio interference. While all wireless media is susceptible to intentional or unintentional DoS events (radio jamming, RF interference) the use of centralized radio management WLAN designs provides solutions to address these challenges ¹ .

Table 4 **Comparison of Wired vs. Wireless Support of Application Requirements (continued)**

QoS	Switched Ethernet provides multiple dedicated hardware queues including a strict priority queue for each port providing the ability to support guaranteed QoS policies. Additional per port per VLAN features such as policers provide granular traffic marking and traffic control and protection against misbehaving clients.	Enhancements to WLAN QoS as defined by the 802.11e standards provide the ability for QoS-enabled stations to have the ability to request specific transmission parameters (data rate, jitter, etc.) required to meet strict QoS policy requirements. Currently most WLAN deployments do not support a full 802.11e implementation and can suffer from QoS degradation under very high traffic loads.
Multicast	The extremely low Bit Error Rates (BER) of fiber and copper links combined with dedicated hardware queues ensure an extremely low probability of dropping multicast traffic and thus a very high probability of guaranteed delivery for that multicast traffic. (Multicast traffic is UDP based and does not have inherent re-transmission capabilities. The ability to reliably guarantee delivery of multicast data is dependent on the ability of the network to prevent packet drops.)	Wireless LAN environments experience a higher BER rate than a comparable wired network and do not provide for acknowledged delivery of multicast data between the AP and the client. While WLAN environments support the transmission of multicast traffic they may not meet the needs of high volume loss sensitive multicast applications (Note: 802.11 unicast traffic uses acknowledged transmissions to achieve a similar reliability for unicast traffic to wired networks even with the inherent higher BER.)
Control of Peer to Peer Traffic	Yes, per port ACL's and PVLAN isolation capabilities allow for segmentation of traffic down to the device level	Yes, peer to peer traffic can be blocked by the WLAN system, at the device level.
Authentication	Client authentication (802.1x) is supported in a switched environment but tends to be an add-on technology to a previously existing mature environment and can prove to have a more complicated deployment than in an equivalent wireless environment.	Client authentication protocols are integrated into WLAN standards and incorporated into the existing end station clients. Consistent client authentication policies are the norm for wireless designs.
Location	Location based services are an add-on technology to a previously existing mature environment.	Location based services integrated into current WLAN systems.

1. Layer 3 DoS protection is common to both environments as it is a property of the shared switched infrastructure

It is reasonable to assume that most enterprise campus environments will continue to have variations in business application requirements and will need a combination of both wired and wireless access for years to come. Neither wired nor wireless environments will be solely sufficient to support all business requirements. The challenge for the network designer is to deploy an integrated campus solution that provides the optimal service requirements for all devices based on the principles of the converged network—while still providing a common baseline set of network services and allowing unified operations and management.

Campus Access Services

The ability to negotiate configuration parameters and settings between edge devices and the network infrastructure is a central property of the campus access layer. Traditionally, switching designs, campus or data center, all appeared fundamentally similar. They consisted of basic Ethernet connectivity with the appropriate number of access ports and overall network capacity. As both the data center and the campus environments have evolved, the designs and system requirements have become more specialized

and divergent. One area where this is most apparent is at the access layer. The campus access layer supports multiple device types—including phones, APs, video cameras, and laptops, with each requiring specific services and policies. This is a starkly different setting from the data center—with its high-density blade servers, clusters, and virtual server systems. PoE, client authentication, dynamic QoS, and security services to support an increasingly mobile works force are requirements in the campus access layer that distinguish it from both legacy switching environments and the specialized needs of the data center.

Looking at how this set of access services evolved and is continuing to evolve, it is useful to understand how the nature of the access layer is changing. DHCP was the first mechanism to provide dynamic edge device network configuration and ease the movement of physical devices throughout the network. Dynamic negotiation of the correct IP stack configuration eased moves adds and changes of PCs, printers and other devices. The migration to VoIP and the ability for phones to dynamically negotiate service requirements with the network provided for another major step in this movement to increased user mobility. In addition to leveraging dynamic IP configuration VoIP devices also leveraged dynamic service registration mechanisms (SCCP registration with the Cisco Unified Communications Manager) as well as dynamic network services negotiation. The ability of the phones to negotiate both power requirements, PoE, as well as edge port QoS, topology, and security parameters provided for a fairly sophisticated plug-and-play capability. Cisco Discovery Protocol (CDP) provides the ability for the end device, such as an IP phone, to identify itself to the network and for both the network and the phone to negotiate configuration parameters. An IP phone identifies (via CDP) the VLAN it needs to use for voice traffic and how to remark the CoS bits on the traffic received from the attached PC. Similarly the switch will identify the specific power requirements as well as the correctly set the port QoS configuration based on the presence of a phone on the edge port. Recent enhancements to this dynamic negotiation process—requiring that a phone negotiate both the correct PoE and CDP parameters before being assigned to the voice VLAN—are additional enhancements providing a higher degree of trust and security to this dynamic negotiation process.

Another trend to be aware of is that network discovery and configuration capabilities of CDP are being complemented with the addition of the IEEE LLDP and LLDP-MED protocols. LLDP and LLDP-MED complement and overlap the functionality provided by CDP, but with a number of differences. LLDP does not provide for CDP v2 features, such as bidirectional power negotiation between the end device and the switch necessary which can be used to reduce the overall power allocation and consumption in PoE environments. In most campus networks, it is reasonable to expect that both CDP and LLDP/LLDP-MED capabilities will need to be enabled and supported on all access switch ports. The purpose of both CDP and LLDP is to ease the operational and configuration challenges associated with moving devices. As the end user community becomes increasingly mobile, it will be necessary for some extended period of time to ensure that any device be able to attach to any port in the campus and receive the appropriate network access configuration and services—whether a device supports CDP, LLDP, or both.

The introduction of 802.1X as an authentication method for users and devices is a part of the next phase of dynamic access provisioning. In addition to providing strong authentication, 802.1X can also be used as a means to further configure network services, VLAN assignment, QoS, and port ACL policies. The 802.1X policy assignment is no longer just based on global defaults for each device type, as in the case of an IP phone, but on the specific device or user requirements. Initial deployments of 802.1X into the campus often proved challenging primarily due to the challenges in integrating a 20-plus year legacy of devices and operating systems that exist in the wired environment. Most legacy wired networks had never been designed or deployed with network authentication in mind. Newer features such as MAC Authentication Bypass (MAB), Web Authentication, and the open authentication capabilities being introduced in the Cisco Catalyst switches will provide the ability to address these challenges. Over time, a common authentication system for both wired and wireless—and most importantly devices moving between wired and wireless access domains—will become the common deployment model. This

unification of wired and wireless capabilities will continue as wired access begins the adoption of 802.1ae and 802.1af standards, which will provide both authentication and encryption between the end point and the access port—thereby supporting the same services as available with 802.11i wireless today.

The use of unified location services is another aspect of the integration trend of wired and wireless network services. Location services solve a number of challenges associated with dynamic network environments. The ability to locate a device to aid in problem resolution is more critical when the device has the ability to roam throughout the network with no associated change control process. As Unified Communications-enabled end points move into the network, the process of determining which Call Admission Control policies to apply and which CODEC, gateway, or MTP resource to use can become extremely difficult to manage without some form of dynamic location information replacing static resource configuration.

Application Optimization and Protection Services

The campus network generally provides the highest capacity and the lowest latency of any portion of the enterprise network. Determining whether or not QoS mechanisms—and the traffic prioritization and protection they provide—are needed within the campus has often been an issue of debate for network planners. Experiences with unexpected problems such as Internet worms and other similar events however have convinced most network engineers that it is not safe to assume that mission-critical applications will always receive the service they require without the correct QoS capabilities in place, even with all the capacity in the world.

A number of other factors are also affecting the ability of networks to support enterprise business requirements:

- The introduction of 10 Gigabit links and more advanced TCP flow control algorithms are creating larger traffic bursts and even larger potential speed mismatches between access devices and the core of the network—driving the need for larger queues.
- The growth in peer-to-peer traffic and the overloading of well-known ports with multiple application and traffic types have added another set of challenges. Applications masquerading as web traffic and multiple applications with different service requirements all using the same HTTP ports are both examples of port overloading.
- Traffic flows within the campus are becoming more complex and diverse. The ability to predict the location of congestion points becomes more difficult as data flow patterns are able to migrate while dynamic peer-to-peer sessions come and go from the network.
- The ability to identify the critical vs. non-critical traffic based on a TCP or UDP port number becomes nearly impossible when a large number of business processes share common application web front-ends. It becomes even harder to find unwanted or unknown applications when those applications have been written to use a variety of port numbers and are able to masquerade as HTTP traffic on TCP port 80 while dynamically searching for access through corporate firewalls.

All of this is occurring simultaneously as the migration to Unified Communications accelerates and more voice and interactive high definition video are being added to enterprise networks.

Principles of Campus QoS Design

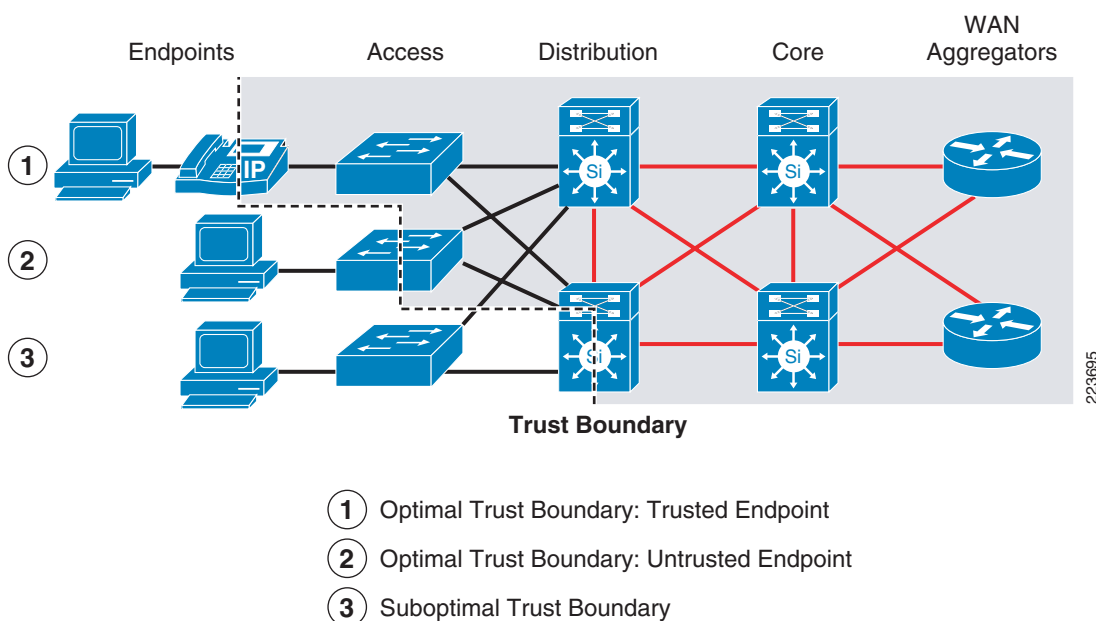
When considering requirements for optimizing and protecting applications and traffic flows in the campus, it is essential to understand what QoS tools are available and how to use. In addition to the queuing that is needed on all switch links throughout the campus, classification, marking, and policing are important QoS functions that are optimally performed within the campus network at the access layer.

Three QoS design principles are important when deploying campus QoS policies:

- Classify and mark applications as close to their sources as technically and administratively feasible. This principle promotes end-to-end Differentiated Services/Per-Hop Behaviors.
- Police unwanted traffic flows as close to their sources as possible. This is especially the case when the unwanted traffic is the result of DoS or worm attacks.
- Always perform QoS functions in hardware rather than software when a choice exists.

Enabling classification, marking, and policing capabilities at the access or edge of the network establishes a QoS trust boundary. The trust boundary is the point in the network where all traffic beyond that point has been correctly identified and marked with the correct Class of Service (CoS)/Differentiated Services Code Point (DSCP) markings. It defines the part of the network in which application flows are protected and those portions in which they are not. Defining the trust boundary as close to the edge of the network as possible means *all* of the application flows—even person-to-person voice calls between colleagues in the same area are protected. See [Figure 23](#).

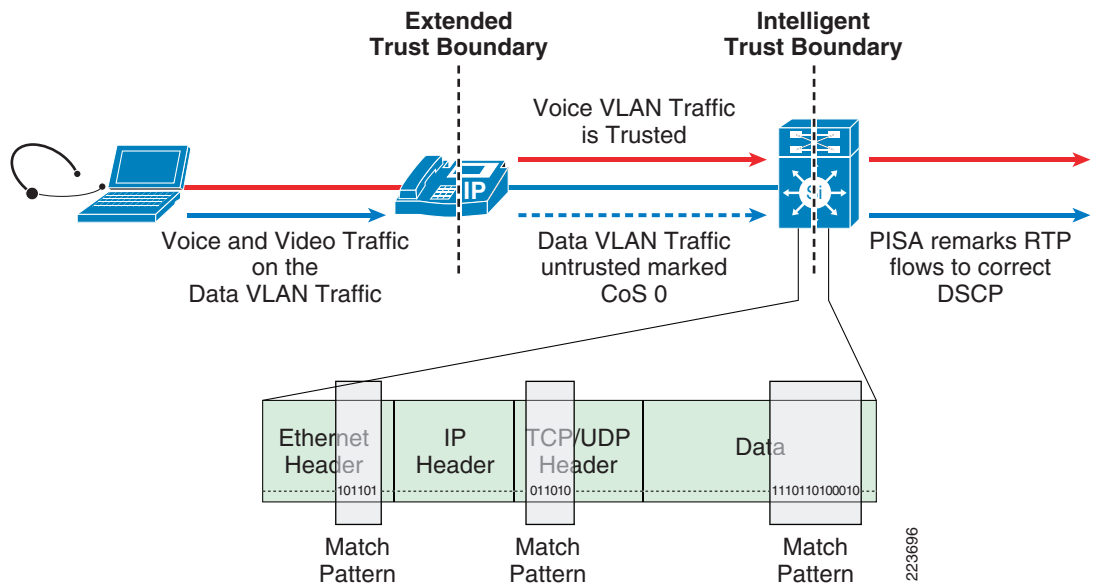
Figure 23 *Campus QoS Trust Boundary Recommendations*



In the current campus QoS design, the access ports of each switch are configured to not trust the QoS markings of any traffic arriving on that port—unless it is on the auxiliary or voice VLAN and the switch has detected that there is a phone (trusted device) on that VLAN. The decision to trust or not trust the endpoints traffic is binary; either the traffic is from the phone and trusted or from any other device and not trusted. This model works well in an environment with dedicated phones, but as the trends in Unified Communications continue and voice/video applications start merging with other PC applications, the need to selectively and intelligently trust certain application flows from the *untrusted* PC is becoming necessary. The use of per VLAN and per port traffic policers is one mechanism that is used to selectively trust traffic in certain port ranges and at certain data rates. Each edge port can be configured to detect traffic within a specific port range and, for all traffic that is less than a defined *normal* rate, mark that traffic with the correct DSCP values. All traffic in excess of this rate is dropped, which provides a safety mechanism to protect against one application masquerading as another more mission critical one (by using the more important application's port numbers for communication). While this policer-based approach has proven to work well and is still valid for certain environments, the increasingly complex list of applications that share port numbers and applications that might be hijacking other applications trusted port ranges requires that we consider a more sophisticated approach.

Deep Packet Inspection (DPI) or the capability to look into the data payload of an IP packet, and not just use the IP and TCP/UDP header to determine what type of traffic the packet contains, provides a tool to address this problem. A switch equipped with hardware Network Based Application Recognition (NBAR) is able to determine whether a specific UDP flow is truly an RTP stream or some other application-based by examining the RTP header contained within the payload of the packet. See Figure 24.

Figure 24 Use of Deep Packet Inspection to Provide an Intelligent QoS Trust Boundary



The ability to detect and appropriately mark specific application flows at the edge of the network provides for a more granular and accurate QoS trust boundary.

Until recently, it has been recommended that the end devices themselves not to be considered as trusted unless they were strictly managed by the IT operations group. It has always been possible for a user to configure the NIC on their PC to mark all their traffic to any classification. If they marked all traffic to DSCP EF they could effectively hijack network resources reserved for real time applications (such as VoIP), thereby ruining the VoIP service quality throughout the enterprise. The introduction of capabilities in the Cisco Security Agent (CSA) and in Microsoft Vista to provide for centralized control of the QoS classification and marking of application traffic flows is another approach that should allow for a more granular QoS trust policy. It is important to note when considering the overall campus QoS design that the capabilities of the Vista and CSA clients do not provide for policing and other traffic control capabilities offered by the switches. It is still recommended that, in campus environments leveraging the CSA and Vista marking capabilities, the network itself be designed to provide the appropriate traffic identification and policing controls.

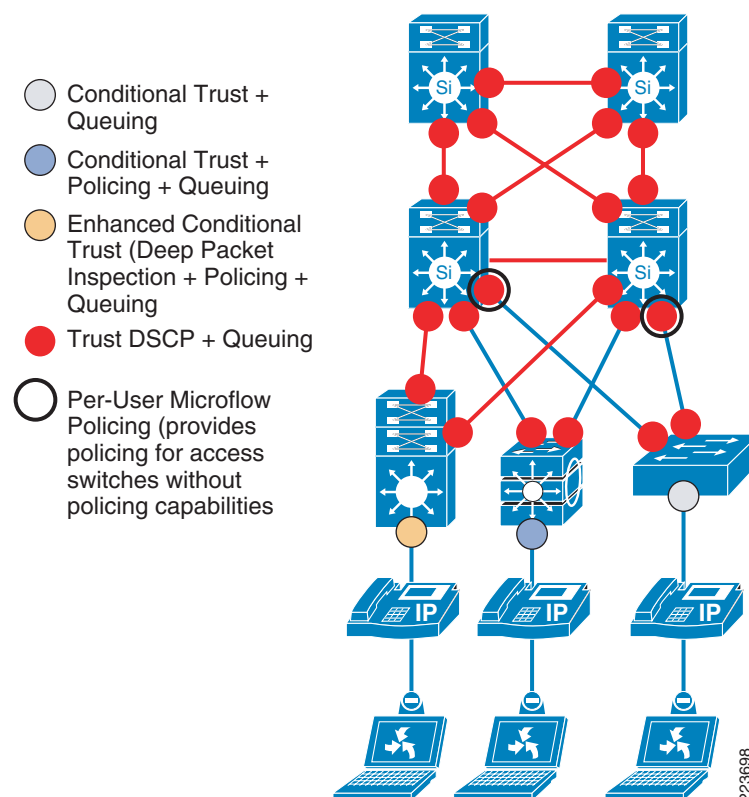


Note

Microsoft has implemented a number of flow control mechanisms into the Vista IP stack that are intended to provide for improved traffic management capabilities. As of the time this document was written, Cisco was still in collaboration with Microsoft to determine the effectiveness and best practices for the use of these new QoS tools. Currently the best practice is still recommended to deploy a traditional trust boundary model complemented by DPI.

The presence of the trust boundary in the campus QoS design provides the foundation for the overall architecture. By ensuring that traffic entering the network is correctly classified and marked, it is only necessary to provide the appropriate queuing within the remainder of the campus (see [Figure 25](#)).

Figure 25 Campus QoS Classification, Marking, Queuing and Policing



Network Resiliency and QoS

The use of QoS in the campus is usually intended to protect certain application traffic flows from periods of congestion. In a campus environment with mission critical applications, the use of QoS tools and design principles provides enhanced resiliency or availability for those mission applications that are explicitly protected based on their CoS/DSCP markings. By enhancing the baseline campus QoS design to include mechanisms such as a *scavenger* queue combined with DPI and edge policing, it is also able to provide for a degree of protection for all of the remaining best effort applications.

The principles behind the use of scavenger classification are fairly simple. There are certain traffic flows in any network that should receive what is termed *less-than-best-effort* service. Applications that do not need to complete in a specific time, such as some types of backups or are non-essential to business processes, can be considered as scavenger traffic. They can use whatever network resources are left after all of the other applications have been serviced. Once a specific traffic flow is determined to fall into this category, all of its packets are marked with DSCP value CS1 to indicate that they are classified as scavenger traffic. Specific queues with a high drop probability are then assigned for the scavenger traffic that provide a throttling mechanism in the event that the scavenger traffic begins to compete with the best-effort flows.

Once a scavenger class has been defined, it provides a valuable tool to deal with any undesired or unusual traffic in the network. By using NBAR (deep packet inspection), it is possible to determine that there are undesired applications on the network and either drop that traffic or mark it as scavenger—depending on the type of traffic and the network policy. By implementing an ingress policer on access ports in the campus, it is also possible to determine whether any device or application begins to transmit at abnormally high data rates. Traffic that exceeds a normal or approved threshold for an extended period of time can also be classified as scavenger.

Having a QoS design and policy that identifies unwelcome or unusual traffic as scavenger traffic provides for additional protection on the fair access to network resources for all traffic—even that marked best effort. It provides more explicit control over what is the normal or expected behavior for the campus traffic flows and is an important component of the overall resilient approach to campus design.

**Note**

For more details on the use of Scavenger QoS and the overall campus QoS design, see the campus QoS design chapter of the *Enterprise QoS Solution Reference Network Design Guide Version 3.3* which can be found on the CCO SRND site, <http://www.cisco.com/go/srnd>.

Virtualization Services

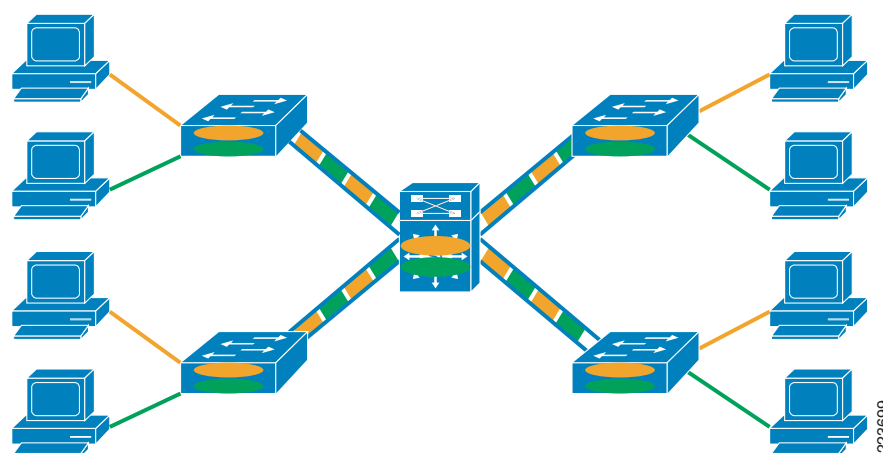
Many enterprises provide network services for departmental networks or business units, hosted vendors, partners, guests. Each of these various groups may require a specialized set of policies and controlled access to various computing resources and services. It is also often the case that certain regulatory or compliance restrictions mandate specific access control, traffic isolation, or traffic path control for certain groups. Some of these groups might exist in the network for long periods of time, such as partners, and others might only require access for the life of a specific project—such as contractors. A network might also find itself having to support a growing number of itinerant guest users. Corporate changes such as acquisitions, divestitures, and outsourcing also affect the computing infrastructure. The manner in which communications and computing are intertwined into the enterprise business processes means that any change in the structure of the organization is immediately reflected in the needs of the campus and the network as a whole. The requirement for a campus network to rapidly respond to these sudden changes in business policy demands a design with a high degree of inherent flexibility.

Virtualization—the ability to allocate physical resources in a logical fashion (one physical device shared between multiple groups or multiple devices operated as a single logical device)—provides the ability to design in a high degree of flexibility into the campus architecture. Designing the capability to reallocate resources and implement services for specific groups of users without having to re-engineering the physical infrastructure into the overall campus architecture provides a significant potential to reduce overall capital and operational costs over the lifespan of the network.

Campus Virtualization Mechanisms

Virtualization capabilities are not new to the campus architecture. The introduction of Virtual LANs (VLANs) provided the first virtualization capabilities in the campus. See [Figure 26](#). The ability to have one device, a switch, replace multiple hubs and bridges while providing distinct forwarding planes for each group of users was a major change to the campus design.

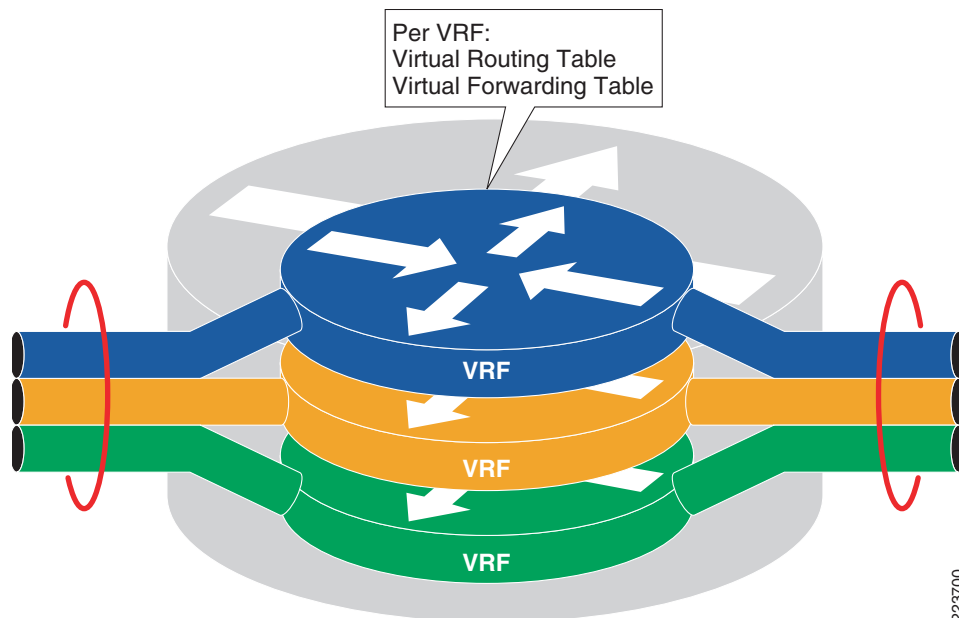
Figure 26 Virtual LAN (Campus Virtualization)



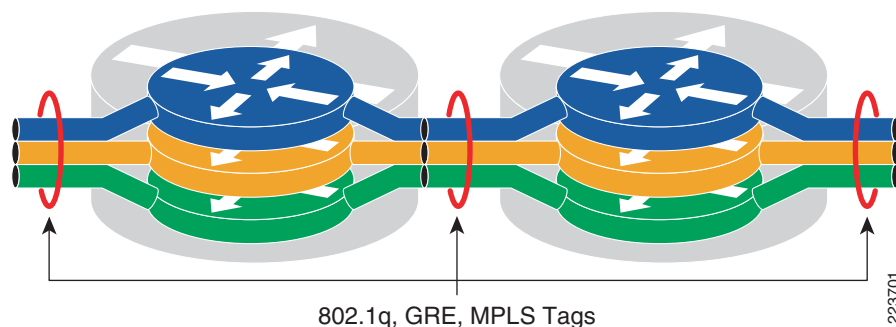
The use of a switched VLAN-based design has provided for a number of advantages, increased capacity, isolation and manageability. However, it is the flexibility that VLANs offer that has had the largest impact on campus designs. The ability to dynamically reconfigure the network, add new subnets or business groups, without having to physically replace the network provided huge cost and operational benefits. Today's modern campus networking environment exists largely due to the capabilities that VLAN virtualization provided.

While VLANs provide some flexibility in dynamically segmenting groups of devices, VLANs do have some limitations. As a Layer-2 virtualization technique, VLANs are bound by the rules of Layer-2 network design. In the structured hierarchical campus design do not have the flexibility to span large domains. The use of Virtualized Routing and Forwarding (VRF) with GRE, 802.1q and MPLS tagging to create Virtual Private Networks (VPN) in the campus provides one approach to extending the configuration flexibility offered by VLANs across the entire campus and if required through the entire network. See [Figure 27](#).

VRFs provide the ability to have separate routing and forwarding instances inside one physical switch. Each VRF has its own Layer-3 forwarding table. Any device in a specific VRF can be Layer-3 directly switched (in other words, routed) to another device in the same VRF, but cannot directly reach one in another VRF. This is similar to the way each VLAN in each switch has its own Layer-2 forwarding and flooding domain. Any device in a VLAN can directly reach another device at Layer-2 in the same VLAN, but not a device in another VLAN unless it is forwarded by a Layer-3 router.

Figure 27 Virtual Routing and Forwarding (VRF)

Just as with a VLAN based network using 802.1q trunks to extend the VLAN between switches, a VRF based design uses 802.1q trunks, GRE tunnels, or MPLS tags to extend and tie the VRFs together. See [Figure 28](#).

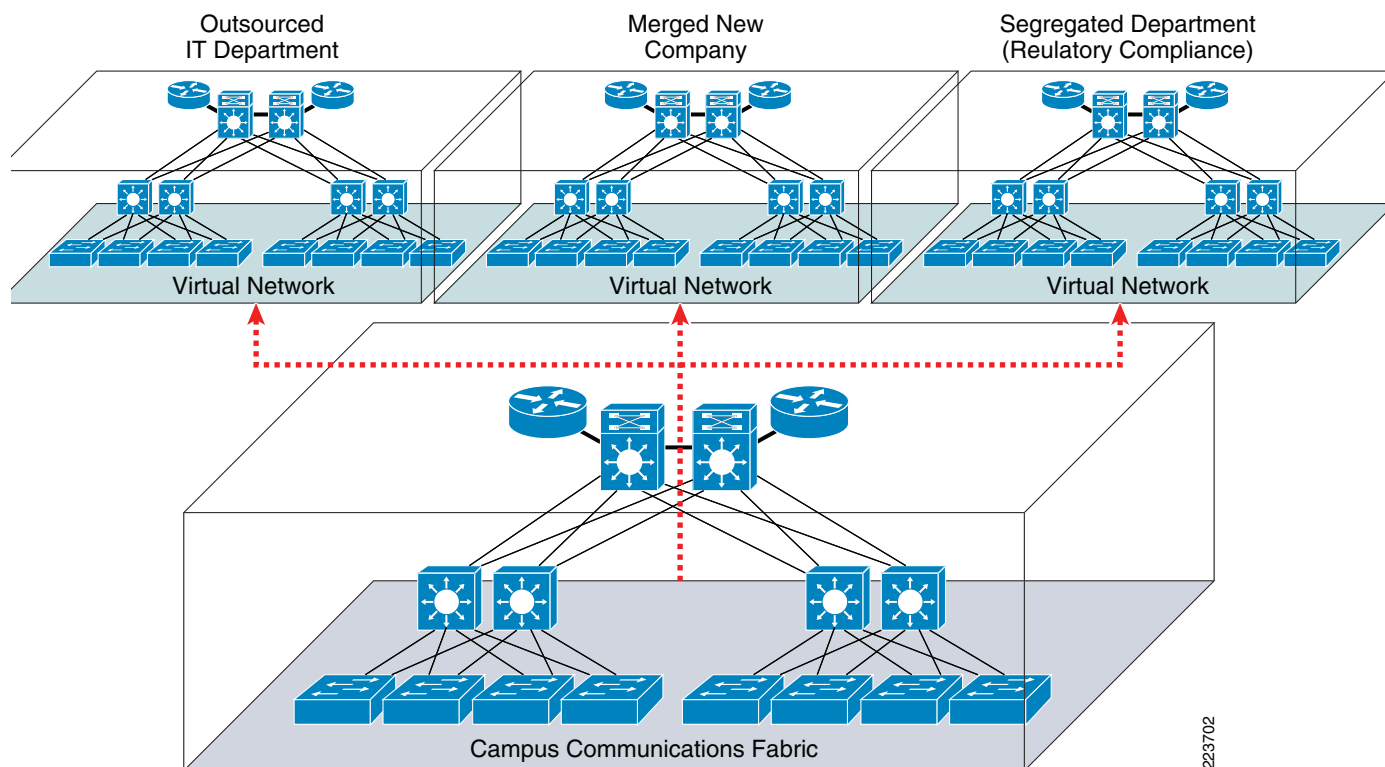
Figure 28 Link Virtualization Options

Any or all of these three link virtualization mechanisms can be used in VRF-based Layer-3 forwarding virtualization in the end-to-end design. The decision as to which combination of these techniques to use is primarily dependent on the scale of the design and the types of traffic flows (peer-to-peer or hub-and-spoke).

Network Virtualization

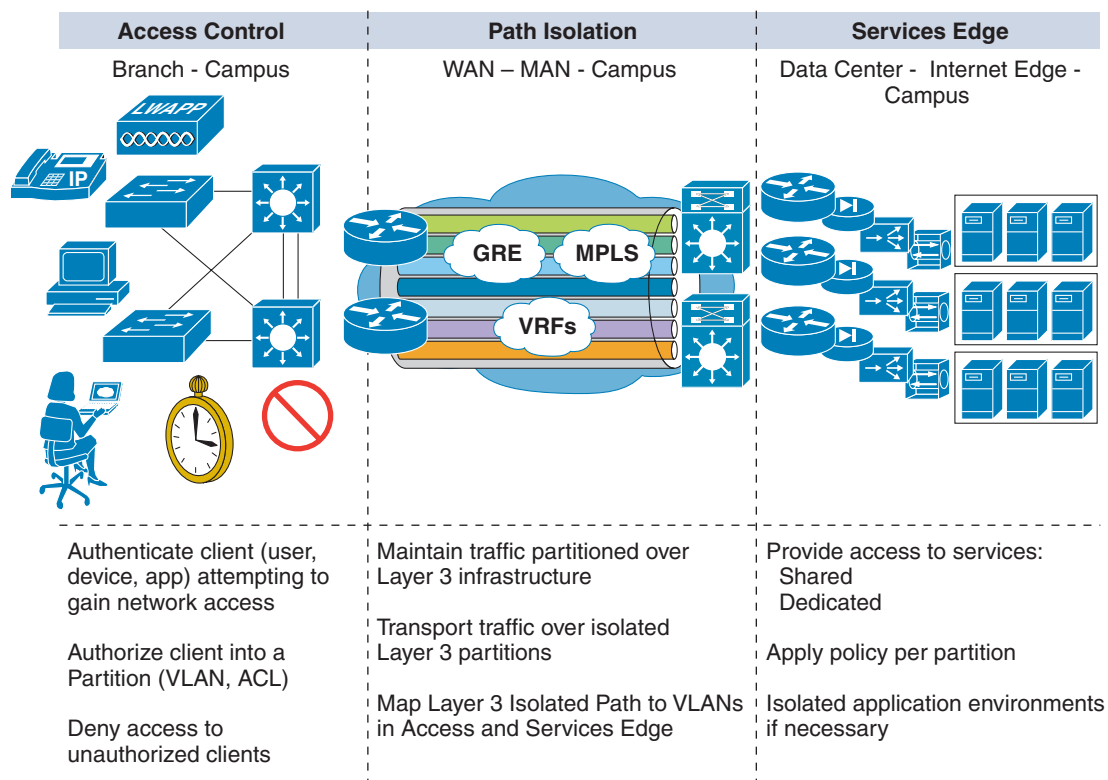
Network Virtualization is best described as the ability to leverage a single physical infrastructure and provide multiple virtual networks each with a distinct set of access policies and yet support all of the security, QoS, Unified Communication services available in a dedicated physical network. Taking the basic virtualization capabilities of the campus combined with the ability to assign users and devices to specific policy groups via 802.1X provides for flexibility in the overall campus architecture. As illustrated in [Figure 29](#), a single physical campus can allow for the allocation of multiple separate logical networks when built with the necessary capabilities.

Figure 29 Example of the Many-to-One Mapping of Virtual to Physical Networks



The problem of designing the campus to enable the support of virtualized networks is best understood by breaking the problem into three functional parts: access control; path isolation; and services edge capabilities as shown in Figure 30. Each of these three parts is in turn built using many individual features—all designed to interoperate and produce the end-to-end virtualized networking solution.

Figure 30 Functional Elements Needed in Virtualized Campus Networks



223703

Enabling access control requires that some form of policy and group assignment be performed at the edge of the network. This can be done dynamically via 802.1X, MAB, Web-Auth, or the NAC appliance. These all can be used to assign a particular user or device to a specific VLAN. It can also be accomplished statically via manual configuration that assigns specific ports to specific VLANs (and specific virtual networks). Path isolation can be accomplished via any combination of the virtual forwarding and link mechanisms. One example is VRF-Lite using VRFs combined with 802.1q trunks, as describe in the preceding description. The services edge policies can be implemented in the data center or in larger networks locally in the campus services block module.



Note

For specific details on how each of these three functional areas are implemented in a campus design, see the Network Virtualization section on the SRND page at <http://www.cisco.com/go/srnd>.

Security Services

Security services are an integral part of any network design. The interconnectedness of networks, the increasing use of mobile devices and the change of the mindset of the hacker community—from one where technical pride motivated most attacks to one where financial interests are a primary motivator—have all been responsible for the continuing increase in the security risks associated with our network infrastructures.

Many of the campus *security* features have already been discussed in some form in the various preceding sections. Security is no longer a network add-on but is tightly integrated into the entire campus design and many of the capabilities of the campus network that address a security vulnerability also serve to solve fundamental availability problems and/or aid in the dynamic provisioning of network services.

Within the networked environment today, there are a wide variety of attack vectors and types—ranging from the simple data sniffing to sophisticated *botnet* environments leveraging complex distributed control systems. All of these various security attacks fall within six fundamental classes of security threats that the campus design must consider:

- Reconnaissance attacks
- Denial of service/distributed denial of service attacks
- Eavesdropping attacks
- Collateral damage
- Unauthorized access attacks
- Unauthorized use of assets, resources, or information

Addressing these threats requires an approach that leverages both prevention and detection techniques to address the root cause attack vectors or vulnerabilities that security hacks use—as well as provide for rapid response in the event of an outbreak or attack. Combining tools within the switching fabric with external monitoring and prevention capabilities will be necessary to address the overall problem.

The security architecture for the campus can be broken down into three basic parts: infrastructure; perimeter and endpoint security; and protection. These are addressed in the sections that follow.

Infrastructure Security

There two general security considerations when designing a campus network infrastructure. First, the infrastructure must be protected from intentional or accidental attack—ensuring the availability of the network and network services. Secondly, the infrastructure must provide information about the state of the network in order to aid in detection of an ongoing attack.

Infrastructure Protection

The security design must provide protection for three basic elements of the infrastructure: devices (switches); links; and, the control plane.

Protecting the Network Devices

Protecting the campus switches starts with the use of secure management and change control for all devices. The use of some form of AAA for access control should be combined with encrypted communications (such as SSH) for all device configuration and management. The preferred AAA methods are RADIUS or TACACS+; these should be configured to support command authorization and full accounting. As an additional step, each device should be configured to minimize the possibility of any attacker gaining access or compromising the switch itself. This protection is accomplished using the Cisco IOS AutoSecure feature. AutoSecure is a Cisco IOS system macro that updates each switch's

security configuration to bring it inline with the Cisco-recommended security best practices. While the use of the AutoSecure feature can greatly ease the process of protecting all the devices in the network, it is recommended that a network security policy be developed and that a regular audit process be implemented to ensure the compliance of all network devices.

Protect the Links

Protecting the inter-switch links from security threats is largely accomplished through the implementation of the campus QoS design discussed in the [Application Optimization and Protection Services, page 38](#). Having the appropriate trust boundary and queuing policies—complemented with the use of scavenger tools in the overall design—will aid in protecting the link capacity within the trusted area (inside the QoS trust boundary) of the network from direct attack. Areas outside of the QoS trust boundary will require additional mechanisms, such as the Cisco DDoS Guard, deployed to address the problems of link saturation by malicious attack.

Protect the Control Plane

Protecting the control plane involves both hardening the system CPU from overload conditions and securing the control plane protocols. The use of MD5-based authentication and explicitly disabling any control protocol on any interface where it is not specifically required, together provide the first level of protection by securing the control plane protocols. Once these exposures have been closed, the next problem is protecting the switch's CPU from other vulnerabilities. If the CPU of the switch can be attacked and overloaded—either intentionally or unintentionally—the control plane is also vulnerable. If the switch is unable to process routing, spanning tree, or any other control packets, the network is vulnerable and its availability is potentially compromised. As discussed in the [Tools and Approaches for Campus High Availability, page 30](#), this type of problem is best addressed with CPU rate limiting tools (either hardware rate limiters or hardware queuing algorithms) combined with an intelligent Control Plane Policing (CoPP) mechanism. Security, QoS, and availability design overlap here as we need to use QoS tools to address a potential security problem that is directly aimed at the availability of the network.

Infrastructure Telemetry and Monitoring

Without the ability to monitor and observe what is happening in the network, it can be extremely difficult to detect the presence of unauthorized devices or malicious traffic flows. The following mechanisms can be used to provide the necessary telemetry data required to detect and observe any anomalous or malicious activities:

- *NetFlow*—Provides the ability to track each data flow that appears in the network.
- *Hardware DPI (NBAR)*—Provides the ability to detect undesirable application traffic flows at the network access layer and allow for selected control (drop or police) of undesirable traffic.
- *Syslog*—Provides the ability to track system events.

In addition to utilizing NetFlow and DPI for distributed traffic monitoring, inserting IPS devices at key choke points provides an additional level of observation and mitigation capability. While NetFlow provides for a very scalable mechanism to detect and find anomalous traffic flows, IPS along with NBAR based DPI can provide visibility into the content of individual packets. All three of these telemetry mechanisms must be supported by the appropriate backend monitoring systems. Tools, such as the Cisco MARS, should be leveraged to provide a consolidated view of gathered data to allow for a more accurate overall view of any security outbreaks.



Note

An upcoming campus design chapter will document the detailed best practices for implementing campus infrastructure security and hardening as outlined above.

Perimeter Access Control and Edge Security

Just as a firewall or external security router provides security and policy control at the external perimeter of the enterprise network, the campus access layer functions as an internal network perimeter. The network should be able to provide the reassurance that the client connecting at the internal perimeter is indeed a known and trusted client (or at least meets the minimal requirements to be safely allowed to connect at this point in the network). Trust and identity features should be deployed at these internal perimeters in the form of authentication mechanisms such as IBNS (802.1X) or Network Admission Control (NAC). This allows the prevention of unauthorized access and/or the ability to introduce compliance and risk management at connection time. Preventing unauthorized access also mitigates the threat of compromise to additional assets in the network.

In addition to ensuring the authentication and compliance of devices attaching to the network, the access layer should also be configured to provide protection against a number of Layer-2 *man-in-the-middle* (MiM) attacks. Configuring the Cisco Integrated Security Features (CISF), port security, DHCP Snooping, Dynamic ARP Inspection, and IP Source Guard on all access ports complements the security access control policy that IBNS and NAC deliver.

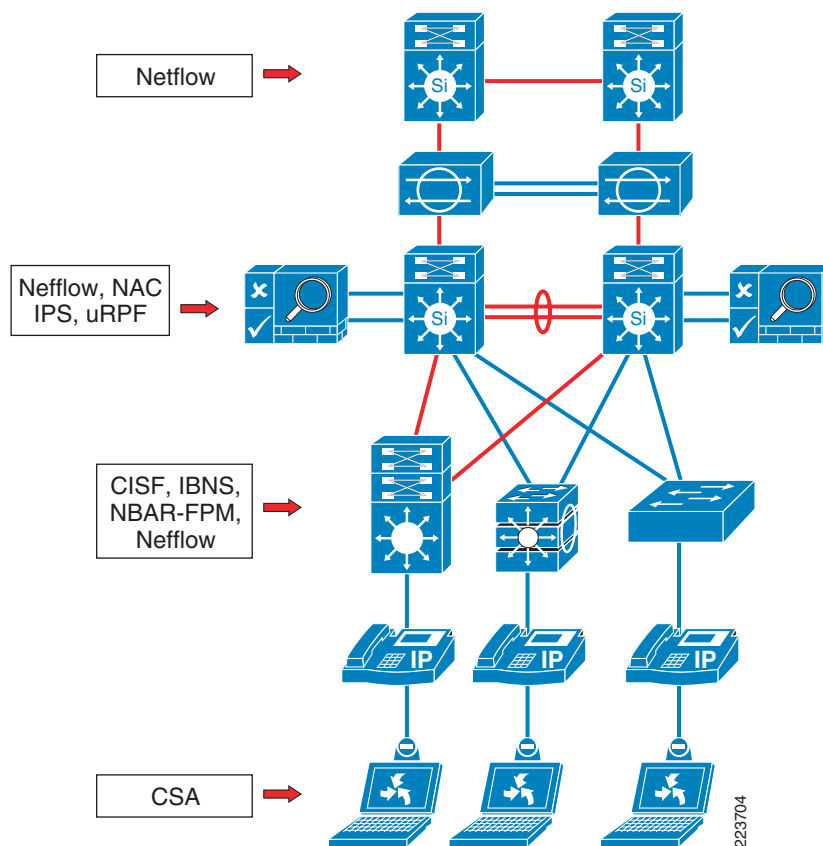
Endpoint Security

The campus security architecture should be extended to include the client itself. Endpoints, such as laptops, are the most vulnerable and most desirable targets for attack. They contain important data and, when compromised, can also serve as a launching points for other attacks against the internal network. The growing threat of *bots* is just the latest in a long line of endpoint vulnerabilities that can threaten the enterprise business.

The installation of client applications, such as Cisco Security Agent (CSA), is an important step towards completing the end-to-end security architecture—along with NAC and IBNS client software on the endpoints that participate with the rest of the integrated network security elements. It is one part of the effort to aid the complex operations of application level security by leveraging the networks integrated security services.

Distributed Security—Defense in Depth

Perhaps the largest security challenge facing the enterprise today is one of scale. The problem of how to detect, prevent, and mitigate against the growing number of security threats requires an approach that leverages a set of security tools that scale proportionally with the size of the network. One approach to this problem of scale is to distribute the security services into the switching fabric itself. An example of this approach is illustrated in [Figure 31](#). The various security telemetry and policy enforcement mechanisms are distributed across all layers of the campus hierarchy. As the network grows in the distributed model, the security services grow proportionately with the switching capacity.

Figure 31 **Distributed Security Services**

In addition to providing a scalable approach to campus security, the distributed model tends to re-enforce a *depth-in-defense* stance. By integrating security functions at all levels of the network, it becomes easier to provide for redundant security monitoring and enforcement mechanisms.

Operational and Management Services

Ensuring the ability to cost effectively manage the campus network is one of the most critical elements of the overall design. As the investment cycle for campus networks lengthens, the operational network costs (OPEX) are increasing relative to the original capital expenditures (CAPEX). Devices remain in service longer and the percentage of overall cost associated with the long-term operation of each device is growing relative to its original capital cost. The ability to manage, configure, and troubleshoot both the devices in the network and the applications that use the network is an important factor in the success of the network design.

The FCAPS framework defines five network management categories: Fault; configuration; accounting, performance; and, security. A full discussion of network management and a comprehensive examination of each of these areas is outside of the scope of this document; however, understanding the principles of campus design and switch capabilities within the overall management framework is essential. Each is described briefly in the sections that follow.

Fault Management

One of the primary objectives of the overall campus design is to minimize the impact of any fault on the network applications and services. The redundancy and resiliency built into the design are intended to prevent failures (faults) from impacting the availability of the campus. Failures will still occur however and having the capabilities in place to detect and react to failures as well as provide enough information to conduct a post mortem analysis of problems are necessary aspects of sound operational processes. Fault management process can be broken down into three stages or aspects, proactive, reactive and post mortem analysis.

Proactive Fault Management

Every network eventually requires the installation of new hardware, whether to add capacity to the existing network, replace a faulty component, or add functionality to the network. The ability to proactively test this new hardware and ensure that it is functioning correctly prior to installation can help avoid any further service interruptions once equipment is installed in the network. While all vendors extensively test and certify that equipment is working correctly before it is shipped to a customer, many things can happen to a piece of equipment before it is finally installed into the production network. Equipment can be damaged during shipping or damaged during installation (static discharge can damage electronic components if systems are not installed using the correct procedures). While care is taken to ensure none of these events occur, having the capability to run extensive diagnostics to detect any failed components prior to any production cutover can avoid potential production problems from occurring later.

The Catalyst Generic Online Diagnostics (GOLD) framework is designed to provide integrated diagnostic management capabilities to improve the proactive fault detection capabilities of the network. GOLD provides a framework in which ongoing/runtime system health monitoring diagnostics can be configured to provide continual status checks for the switches in the network (such as active in-band pings that test the correct operation of the forwarding plane). GOLD also provides the capability to run (or schedule) potentially intrusive on-demand diagnostics. These diagnostics can aid in troubleshooting suspected hardware problems and provide the ability to proactively test new hardware before production cutovers.

**Note**

For more information on GOLD, refer to the following URL:

http://www.cisco.com/en/US/partner/products/ps7081/products_white_paper0900aecd801e659f.shtml

Reactive Fault Management

One of the central objectives for any campus design is to ensure that the network recovers intelligently from any failure event. The various control protocols (such as EIGRP or OSPF) all provide the capability to configure specific responses to failure events. However, in some cases the standard control protocol capabilities are not sufficient and the design might require an additional level of customization as a part of the recovery process. Traditional approaches to adding this customized behavior often involve the use of centralized monitoring systems to trap events and run scripts to take a specific action for each type of event. Providing additional distributed intelligence in the switching fabric can complement and/or simplify these operational processes. Tools, such as the Cisco IOS Embedded Event Manager (EEM), provide the capability to distribute the scripts to switches in the network—rather than running all scripts centrally in a single server. Distributing the scripting intelligence into the campus network itself leverages the distributed processing capacity and direct fault monitoring capabilities of the switches. Capabilities, such as Enhanced Object Tracking (EOT), also provide an additional level of configurable

intelligence to the network recovery mechanisms. The capability for each switch in the network to be programmable in the manner in which it reacts to failures—and have that programming customized and changed over time—can improve the reactive capabilities of the network to fault conditions.

Post Mortem Analysis Capabilities

It is important for the network to recover from the failure when a failure occurs. It is also important in the drive towards maintaining a high level of overall network availability that the operations teams be able to understand what went wrong. Having a centralized record of network events (via SNMP and syslog data), provides for the first level or network topology view of post mortem diagnostic information. In order to provide a more detailed view of specific failure events within the individual devices, it is necessary for the devices themselves to gather and store more detailed diagnostic data. Since centralized management systems are unable to gather data from a device that is no longer fully operational (if that part of the network is down you can not gather data via the network), it is important to have a local store of event information. Some mechanisms—such as the Catalyst System Event Archive (SEA)—can store a record of all local system events in non-volatile storage across reboots. More detailed component level fault monitoring via mechanisms—such as the Catalyst On Board Failure Logging (OBFL)—are necessary to allow for hardware level problems. OBFL acts as a black box recorder for line cards and switches. It records operating temperatures, hardware uptime, interrupts, and other important events and messages that can assist with diagnosing problems with hardware cards (or modules) installed in a Cisco router or switch. Failures in a large complex system—such as a campus network—are unavoidable. Having the capabilities designed into the network to support a post mortem problem analysis process is highly valuable to any enterprise aiming for a high *number of nines* of availability.

Accounting and Performance

Accounting and performance are two aspects of the FCAPS model that are primarily concerned with the monitoring of capacity and the billing for the use of the network. Enterprise environments are not usually as concerned with the accounting aspects of the FCAPS model because they usually do not implement complex usage billing systems. However, enterprises do require the ability to observe the impact of the network on application traffic and end-systems performance. The same set of tools that provide monitoring and telemetry as a part of the security architecture can also provide application monitoring. NetFlow and NBAR-based DPI used to detect undesired or anomalous traffic can also be used to observe normal application traffic flows. Increases in the volume of application traffic—or the detection of new application traffic patterns that might require network upgrade or design changes—can be tracked via NetFlow. Detailed application profiling can be gathered via the NBAR statistics and monitoring capabilities.

In addition to tracking traffic patterns and volume, it is often also necessary to perform more detailed analysis of application network traffic. Distributed network analysis tools (such as packet capture and RMON probes) are often very useful elements to include in the overall campus design. These provide the ability to collect packet traces remotely and view them at a central management console. While distributed packet analyzers are powerful tools, it is not always possible to connect one to every switch in the network. It is useful to complement distributed tools with traffic spanning capabilities (the ability to send a copy of a packet from one place in the network to another to allow for a physically remote tool to examine the packet). The basic port spanning capability of each switch should be complemented by the use of remote span (RSPAN) and Encapsulated RSPAN (ERSPAN) to provide this capability. Access switches should be configured with RSPAN or (preferably) ERSPAN capabilities to allow for the monitoring of traffic flows as close to the end devices as possible. ERSPAN is the preferred solution because it allows for the spanned traffic to be carried over multiple Layer-3 hops allowing for the consolidation of traffic analysis tools in fewer locations.

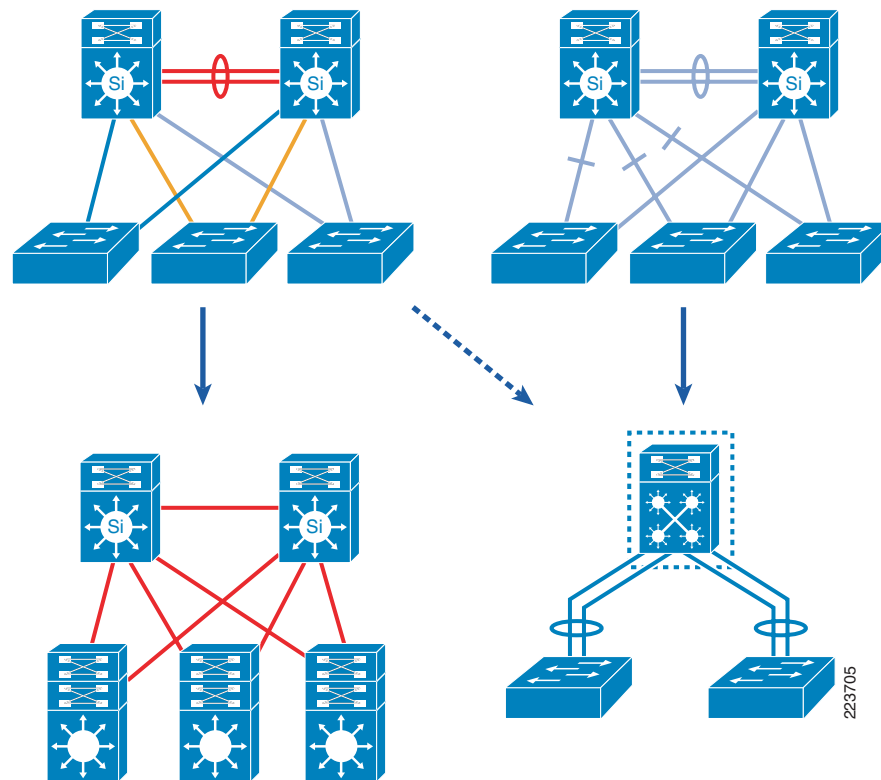
Configuration and Security

The configuration and security of the network devices has been discussed above in the section on security services. The design guidelines described there are intended to meet the needs of the FCAPS model as well as providing a more comprehensive end-to-end campus security. See the “[Security Services](#)” section on page 47 for more information.

Evolution of the Campus Architecture

The campus network architecture is evolving in response to a combination of new business requirements, technology changes, and a growing set of end user expectations. The migration from the more than 10-year-old multi-tier distribution block design to one of the newer routed access-based or virtual switch-based distribution block design options is occurring in response to changing business requirements. See Figure 32. While the traditional multi-tier design still provides a viable option for certain campus environments, increased availability, faster convergence, better utilization of network capacity, and simplified operational requirements offered by the new designs are combining to motivate a change in foundational architectures.

Figure 32 Evolution of the Campus Distribution Block Design



Evolutionary changes are occurring within the campus architecture. One example is the migration from a traditional Layer-2 access network design (with its requirement to span VLANs and subnets across multiple access switches) to a virtual switch-based design. Another is the movement from a design with subnets contained within a single access switch to the routed-access design.

As discussed throughout this document, another major evolutionary change to the campus architecture is the introduction of additional services, including the following:

- Non-stop, high-availability services
- Access and mobility services
- Application optimization and protection services
- Virtualization services
- Security services
- Operational and management services

The motivation for introducing these capabilities to the campus design have been described throughout this document. The increase in security risks, need for more flexible infrastructure, change in application data flows, and SLA requirements have all driven the need for a more capable architecture. However, implementing the increasingly complex set of business-driven capabilities and services in the campus architecture can be a challenge, if done in a piece meal fashion. As outlined in this document, any successful architecture must be based on a foundation of solid design theory and principles. For any enterprise business involved in the design and/or operation of a campus network, we recommend the adoption of an integrated approach—based on solid systems design principles. The *Cisco ESE Campus Design Guide*, which includes this overview discussion and a series of subsequent detailed design chapters, is specifically intended to assist the engineering and operations teams develop a systems-based campus design that will provide the balance of availability, security, flexibility, and operability required to meet current and future business and technological needs.