

Handouts for CS407: Routing and Switching

Topic 1: Routing and Switching - Course Information

This topic provides motivation to study this course.

The focus of this course is **education** and to help you develop a real **understanding** of networking, not merely memorizing facts and commands.

We have divided this course into three parts:

1. Network Fundamentals

In this part, you will learn both the practical and conceptual skills that build the foundation for understanding basic networking.

2. Routing Protocols and Concepts

The focus of this part is on learning the architecture, components, and operations of routers, and explains the principles of routing and routing protocols.

3. Switching Techniques

The focus of this part is on learning the architecture, components, and operations of a converged switched network.

By the end of this course, you will be able to build simple LANs, perform basic configurations for routers and switches, and implement IP addressing schemes.

Figures and Materials used have been adapted from " <http://www.cabrillo.edu/~rgraziani/> " unless stated otherwise.

Topic 2: Components of a Network

Among all of the essentials for human existence, the need to interact with others ranks just below our need to sustain life. Communication is almost as important to us as our reliance on air, water, food, and shelter.

In today's world, through the use of networks, we are connected like never before. People with ideas can communicate instantly with others to make those ideas a reality. News events and discoveries are known worldwide in seconds. Individuals can even connect and play games with friends separated by oceans and continents.

Networks connect people and promote unregulated communication. Everyone can connect, share, and make a difference.

Networks support the way we learn

Communication, collaboration, and engagement are fundamental building blocks of education. Institutions are continually striving to enhance these processes to maximize the dissemination of knowledge.

Access to high quality instruction is no longer restricted to students living in proximity to where that instruction is being delivered. Online distance learning has removed geographic barriers and improved student opportunity. Online (e-learning) courses can now be delivered over a network. These courses can contain data (text, links), voice, and video available to the students at any time from any place. Online discussion groups and message boards enable a student to collaborate with the instructor, with other students in the class, or even with students across the world. Blended courses can combine instructor-led classes with online courseware to provide the best of both delivery methods.

Network Components: Clients and servers

All computers connected to a network that participate directly in network communication are classified as hosts or end devices. Hosts can send and receive messages on the network. In modern networks, end devices can act as a client, a server, or both. The software installed on the computer determines which role the computer plays.

Servers are hosts that have software installed that enable them to provide information, like email or web pages, to other hosts on the network. Each service requires separate server software. For example, a host requires web server software in order to provide web services to the network.

Clients are computer hosts that have software installed that enable them to request and display the information obtained from the server. An example of client software is a web browser, like Internet Explorer.

Peer to Peer

Client and server software usually runs on separate computers, but it is also possible for one computer to carry out both roles at the same time. In small businesses and homes, many computers function as the servers and clients on the network. This type of network is called a peer-to-peer network.

The simplest peer-to-peer network consists of two directly connected computers using a wired or wireless connection.

Multiple PCs can also be connected to create a larger peer-to-peer network but this requires a network device, such as a hub, to interconnect the computers.

The main disadvantage of a peer-to-peer environment is that the performance of a host can be slowed down if it is acting as both a client and a server at the same time.

The advantages of peer-to-peer networking:

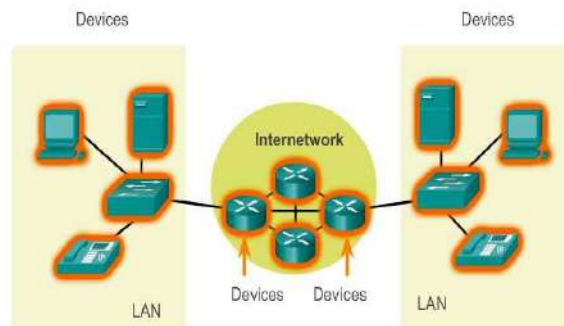
- Easy to set up
- Less complexity
- Lower cost since network devices and dedicated servers may not be required
- Can be used for simple tasks such as transferring files and sharing printers

The disadvantages of peer-to-peer networking:

- No centralized administration
- Not as secure

- Not scalable
- All devices may act as both clients and servers which can slow their performance

Networks Components



The network infrastructure contains three categories of network components:

- Devices
- Media
- Services

End Devices

The network devices that people are most familiar with are called end devices, or hosts. These devices form the interface between users and the underlying communication network.

Some examples of end devices are:

- Computers (work stations, laptops, file servers, web servers)
- Network printers
- VoIP phones
- TelePresence endpoint
- Security cameras
- Mobile handheld devices (such as smartphones, tablets, PDAs, and wireless debit/credit card readers and barcode scanners)

Intermediary Network Devices

Intermediary devices interconnect end devices. These devices provide connectivity and work behind the scenes to ensure that data flows across the network. Intermediary devices connect the individual hosts to the network and can connect multiple individual networks to form an internetwork.

Examples of intermediary network devices are:

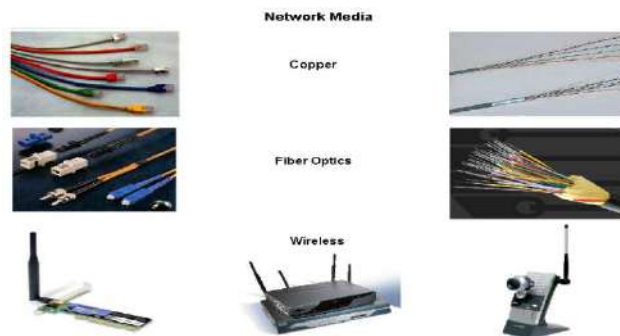
- Network Access (switches and wireless access points)
- Internetworking (routers)
- Security (firewalls)

The management of data as it flows through the network is also a role of the intermediary devices. These devices use the destination host address, in conjunction with information about the network interconnections, to determine the path that messages should take through the network.

Processes running on the intermediary network devices perform these functions:

- Regenerate and retransmit data signals
- Maintain information about what pathways exist through the network and internetwork
- Notify other devices of errors and communication failures
- Direct data along alternate pathways when there is a link failure
- Classify and direct messages according to Quality of Service (QoS) priorities
- Permit or deny the flow of data, based on security settings

Networks Media



Communication across a network is carried on a medium. The medium provides the channel over which the message travels from source to destination.

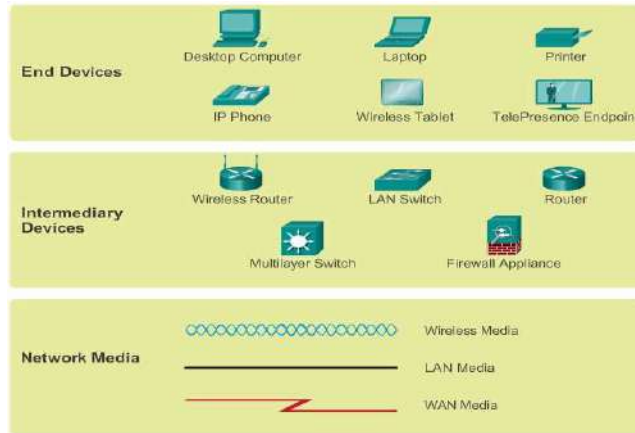
Modern networks primarily use three types of media to interconnect devices and to provide the pathway over which data can be transmitted. As shown in the figure above, these media are:

- Metallic wires within cables
- Glass or plastic fibers (fiber optic cable)
- Wireless transmission

The signal encoding that must occur for the message to be transmitted is different for each media type. Different types of network media have different features and benefits. Not all network media has the same characteristics and is appropriate for the same purpose. The criteria for choosing network media are:

- The distance the media can successfully carry a signal
- The environment in which the media is to be installed
- The amount of data and the speed at which it must be transmitted
- The cost of the media and installation

Network Representation



When conveying complex information such as displaying all the devices and medium in a large internetwork, it is helpful to use visual representations.

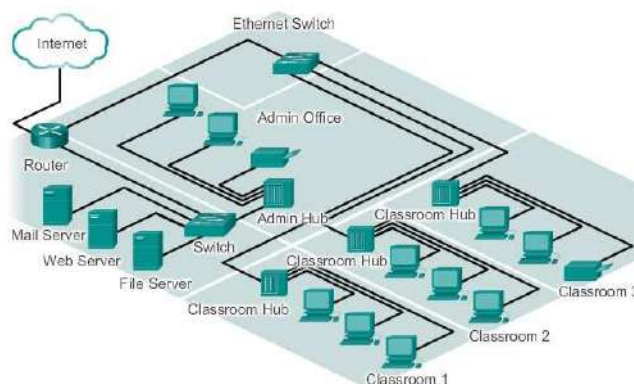
Like any other language, the language of networking uses a common set of symbols to represent the different end devices, network devices, and media, as shown in the figure above. In addition to these representations, specialized terminology is used when discussing how each of these devices and media connect to each other. Important terms to remember are:

Network Interface Card - A NIC, or LAN adapter, provides the physical connection to the network at the PC or other host device. The media connecting the PC to the networking device plugs directly into the NIC.

Physical Port - A connector or outlet on a networking device where the media is connected to a host or other networking device.

Interface - Specialized ports on an internetworking device that connect to individual networks. Because routers are used to interconnect networks, the ports on a router are referred to network interfaces.

Topology Diagrams



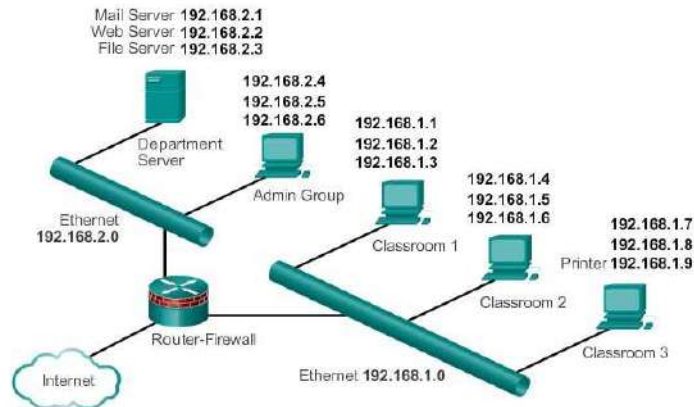
Physical Topology

Topology diagrams are mandatory for anyone working with a network. It provides a visual map of how the network is connected.

There are two types of topology diagrams including:

Physical topology diagrams - Identify the physical location of intermediary devices, configured ports, and cable installation.

Logical topology diagrams - Identify devices, ports, and IP addressing scheme.



Logical Topology

Topic 3: LANs, WANs, and the Internet

Types of Networks

Network infrastructures can vary greatly in terms of:

- Size of the area covered
- Number of users connected
- Number and types of services available

Local Area Network (LAN) - A network infrastructure that provides access to users and end devices in a small geographical area.

Wide Area Network (WAN) - A network infrastructure that provides access to other networks over a wide geographical area.

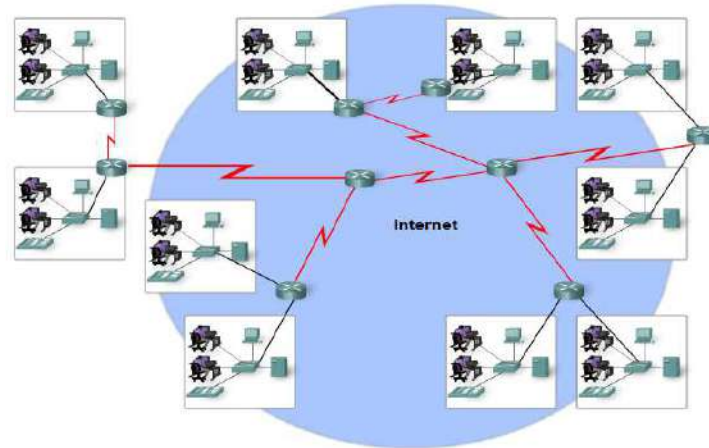
Other types of networks include:

Metropolitan Area Network (MAN) - A network infrastructure that spans a physical area larger than a LAN but smaller than a WAN (e.g., a city). MANs are typically operated by a single entity such as a large organization.

Wireless LAN (WLAN) - Similar to a LAN but wirelessly interconnects users and end points in a small geographical area.

Storage Area Network (SAN) - A network infrastructure designed to support file servers and provide data storage, retrieval, and replication. It involves high-end servers, multiple disk arrays (called blocks), and Fiber Channel interconnection technology.

The Internet - A Network of Networks

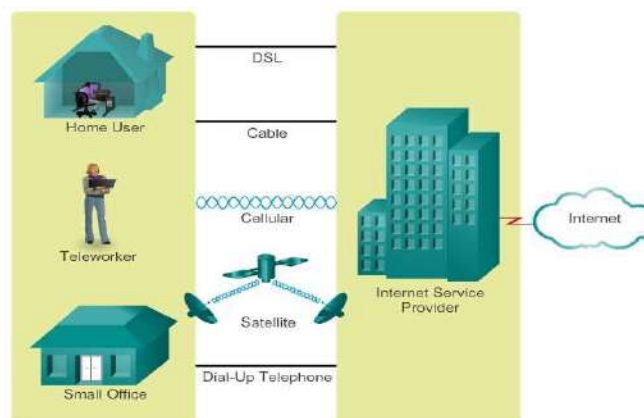


Although there are benefits to using a LAN or WAN, most individuals need to communicate with a resource on another network, outside of the local network within the home, campus, or organization. This is done using the Internet.

As shown in the figure above, the Internet is a worldwide collection of interconnected networks (internetworks or internet for short), cooperating with each other to exchange information using common standards. Through telephone wires, fiber optic cables, wireless transmissions, and satellite links, Internet users can exchange information in a variety of forms.

The Internet is a conglomerate of networks and is not owned by any individual or group. Ensuring effective communication across this diverse infrastructure requires the application of consistent and commonly recognized technologies and standards as well as the cooperation of many network administration agencies. There are organizations that have been developed for the purpose of helping to maintain structure and standardization of Internet protocols and processes. These organizations include the Internet Engineering Task Force (IETF), Internet Corporation for Assigned Names and Numbers (ICANN), and the Internet Architecture Board (IAB), plus many others.

Internet Access Technologies



There are many different ways to connect users and organizations to the Internet. The figure above illustrates common connection options for small office and home office users, which include:

Cable - Typically offered by cable television service providers, the Internet data signal is carried on the same coaxial cable that delivers cable television. It provides a high bandwidth, always on, connection to the Internet. A special cable modem separates the Internet data signal from the other signals carried on the cable and provides an Ethernet connection to a host computer or LAN.

DSL - Provides a high bandwidth, always on, connection to the Internet. It requires a special high-speed modem that separates the DSL signal from the telephone signal and provides an Ethernet connection to a host computer or LAN. DSL runs over a telephone line, with the line split into three channels. One channel is used for voice telephone calls. This channel allows an individual to receive phone calls without disconnecting from the Internet. A second channel is a faster download channel, used to receive information from the Internet. The third channel is used for sending or uploading information. This channel is usually slightly slower than the download channel. The quality and speed of the DSL connection depends mainly on the quality of the phone line and the distance from your phone company's central office. The farther you are from the central office, the slower the connection.

Cellular - Cellular Internet access uses a cell phone network to connect. Wherever you can get a cellular signal, you can get cellular Internet access. Performance will be limited by the capabilities of the phone and the cell tower to which it is connected. The availability of cellular Internet access is a real benefit in those areas that would otherwise have no Internet connectivity at all, or for those constantly on the go.

Satellite - Satellite service is a good option for homes or offices that do not have access to DSL or cable. Satellite dishes require a clear line of sight to the satellite and so might be difficult in heavily wooded areas or places with other overhead obstructions. Speeds will vary depending on the contract, though they are generally good. Equipment and installation costs can be high (although check the provider for special deals), with a moderate monthly fee thereafter. The availability of satellite Internet access is a real benefit in those areas that would otherwise have no Internet connectivity at all.

Dial-up Telephone - An inexpensive option that uses any phone line and a modem. To connect to the ISP, a user calls the ISP access phone number. The low bandwidth provided by a dial-up modem connection is usually not sufficient for large data transfer, although it is useful for mobile access while traveling. A modem dial-up connection should only be considered when higher speed connection options are not available.

Many homes and small offices are more commonly being connected directly with fibre optic cables. This enables an Internet service provider to provide higher bandwidth speeds and support more services such as Internet, phone, and TV.

The choice of connection varies depending on geographical location and service provider availability.

Topic 4: Packet Tracer Basics - Part I

In this topic, we will go through the basics of Packet tracer software. Packet Tracer is a networking learning tool that supports a wide range of physical and logical simulations. It also provides visualization tools to help you understand the internal workings of a network.

Topic 5: Packet Tracer Basics - Part II

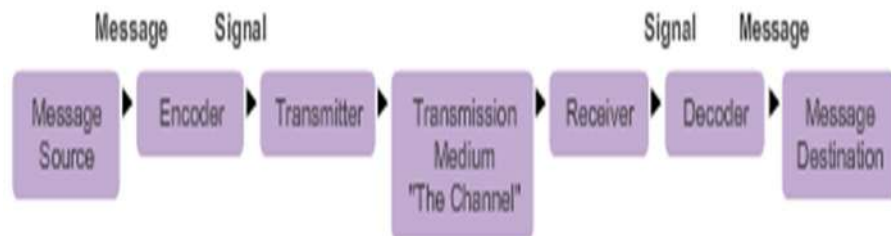
This is the continuation of the previous topic about packet tracer basics.

Topic 6: Rules of Communication

Establishing the Rules

Before communicating with one another, individuals must use established rules or agreements to govern the conversation.

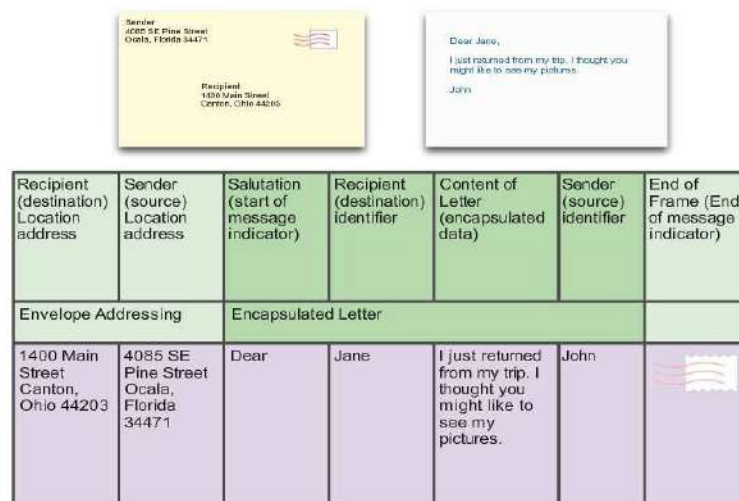
Message Encoding



One of the first steps to sending a message is encoding it. Encoding is the process of converting information into another, acceptable form, for transmission. Decoding reverses this process in order to interpret the information.

Encoding also occurs in computer communication, as shown in Figure above. Encoding between hosts must be in an appropriate form for the medium. Messages sent across the network are first converted into bits by the sending host. Each bit is encoded into a pattern of sounds, light waves, or electrical impulses depending on the network media over which the bits are transmitted. The destination host receives and decodes the signals in order to interpret the message.

Message Formatting and Encapsulation



When a message is sent from source to destination, it must use a specific format or structure. Message formats depend on the type of message and the channel that is used to deliver the message.

A frame acts like an envelope; it provides the address of the intended destination and the address of the source host, as shown in Figure above.

The format and contents of a frame are determined by the type of message being sent and the channel over which it is communicated. Messages that are not correctly formatted are not successfully delivered to or processed by the destination host.

Message Size

Another rule of communication is size. When people communicate with each other, the messages that they send are usually broken into smaller parts or sentences.

Message Timing

Another factor that affects how well a message is received and understood is timing. People use timing to determine when to speak, how fast or slow to talk, and how long to wait for a response. These are the rules of engagement.

Access Method

Access method determines when someone is able to send a message.

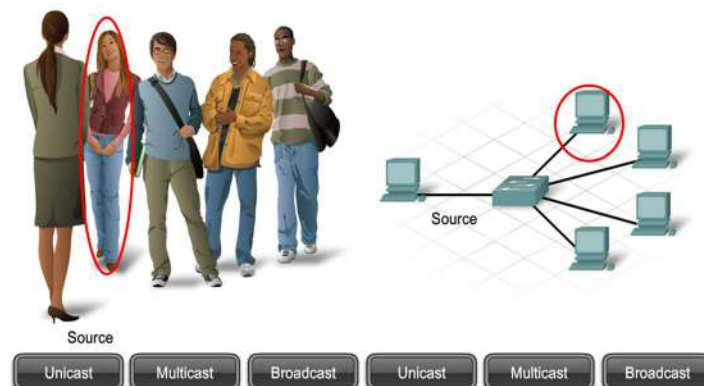
Flow Control

Timing also affects how much information can be sent and the speed that it can be delivered. If one person speaks too quickly, it is difficult for the other person to hear and understand the message. The receiving person must ask the sender to slow down. In network communication, a sending host can transmit messages at a faster rate than the destination host can receive and process. Source and destination hosts use flow control to negotiate correct timing for successful communication.

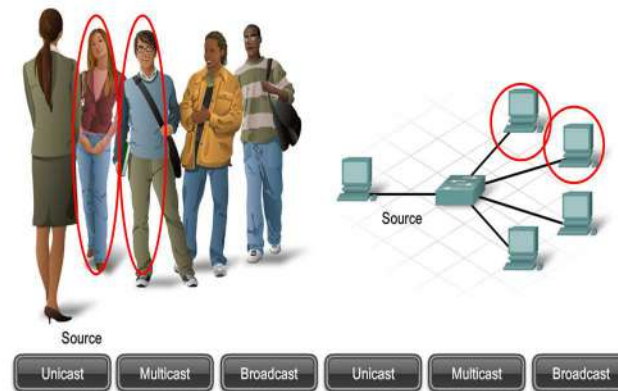
Response Timeout

If a person asks a question and does not hear a response within an acceptable amount of time, the person assumes that no answer is coming and reacts accordingly.

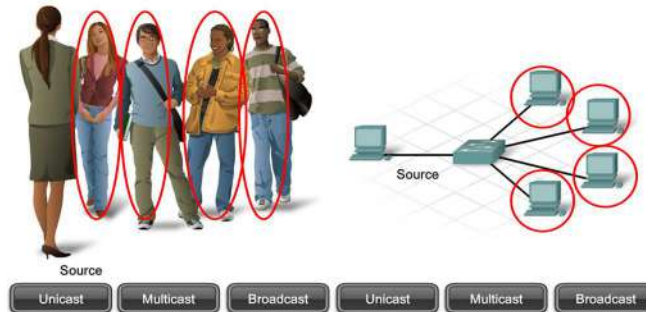
Message Delivery Options - Unicast



Message Delivery Options - Multicast

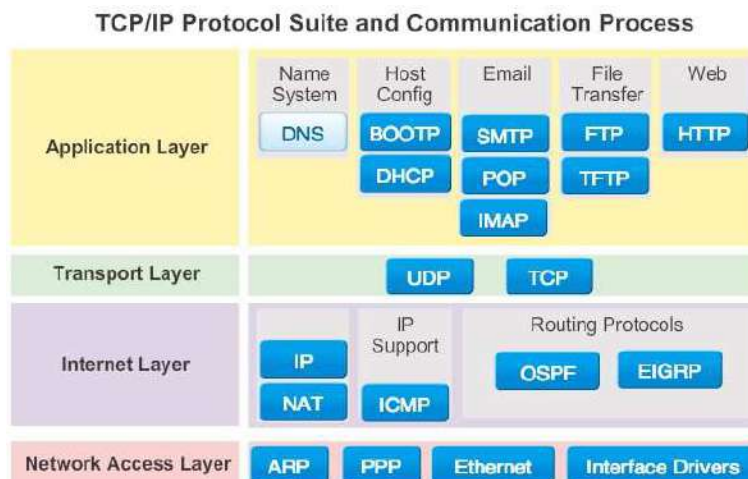


Message Delivery Options - Broadcast



Topic 7: Protocol Suits

TCP/IP Protocol Suite



A protocol suite is a set of protocols that work together to provide comprehensive network communication services. A protocol suite may be specified by a standards organization or developed by a vendor.

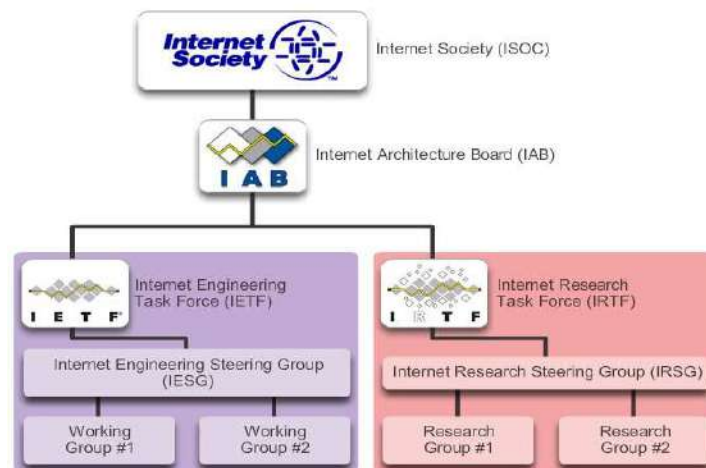
The protocols IP, HTTP, and DHCP are all part of the Internet protocol suite known as Transmission Control Protocol/IP (TCP/IP). The TCP/IP protocol suite is an open standard, meaning these protocols are freely available to the public, and any vendor is able to implement these protocols on their hardware or in their software.

The IP suite is a suite of protocols required for transmitting and receiving information using the Internet. It is commonly known as TCP/IP because the first two networking protocols defined for this standard were TCP and IP. The open standards-based TCP/IP has replaced other vendor proprietary protocol suites, such as Apple's AppleTalk and Novell's Internetwork Packet Exchange/Sequenced Packet Exchange (IPX/SPX).

Today, the suite includes dozens of protocols, as shown in Figure above. They are organized in layers using the TCP/IP protocol model. TCP/IP protocols are included in the internet layer to the application layer when referencing the TCP/IP model. The lower layer protocols in the data link or network access layer are responsible for delivering the IP packet over the physical medium. These lower layer protocols are developed by standards organizations, such as IEEE.

The TCP/IP protocol suite is implemented as a TCP/IP stack on both the sending and receiving hosts to provide end-to-end delivery of applications over a network. The 802.3 or Ethernet protocols are used to transmit the IP packet over the physical medium used by the LAN.

Standard Organizations



- Open standards encourage competition and innovation.
- Guarantee that no single company's product can monopolize the market, or have an unfair advantage over its competition.

Standards organizations include:

The Internet Society (ISOC)

- ISOC facilitates the open development of standards and protocols for the technical infrastructure of the Internet, including the oversight of the Internet Architecture Board (IAB).

The Internet Architecture Board (IAB)

- Responsible for overall management and development of Internet standards.
- Oversight of the architecture for protocols and procedures used by the Internet.
- 13 members, including the chair of the Internet Engineering Task Force (IETF).
- IAB members serve as individuals and not representatives of any company, agency, or other organization.

Internet Engineering Task Force (IETF)

- Mission is to develop, update, and maintain Internet and TCP/IP technologies.
- One of the key responsibilities is to produce Request for Comments (RFC) documents
 - Memorandum describing protocols, processes, and technologies for the Internet.
- The IETF consists of working groups (WGs), the primary mechanism for developing IETF specifications and guidelines.

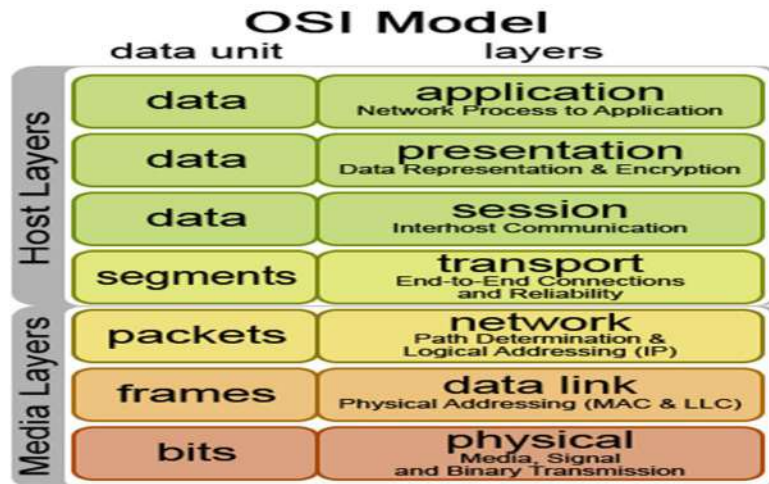
Internet Research Task Force (IRTF)

- Focused on long-term research related to Internet and TCP/IP
- IETF focuses on shorter-term issues of creating standards
- IRTF consists of research groups for long-term development efforts. Including: Anti-Spam Research Group (ASRG), Crypto Forum Research Group (CFRG), Peer-to-Peer Research Group (P2PRG), and Router Research Group (RRG).

Benefits of Layered Model

- Assists in protocol design, because protocols that operate at a specific and a defined interface to the layers above and below.
- Fosters competition
- Prevents technology or capability changes in one layer from affecting other layers above and below.
- Provides a common language to describe networking functions and capabilities.

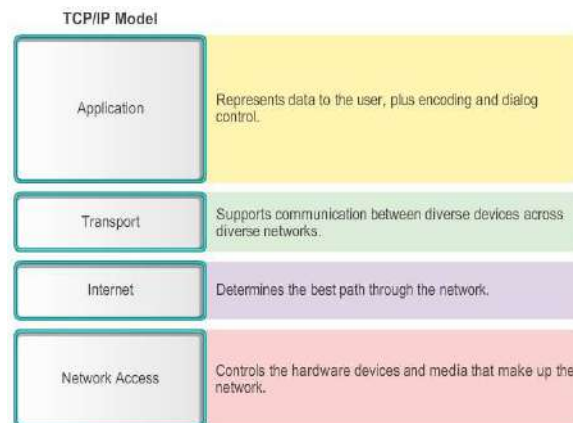
OSI Model



Initially the OSI model was designed by the ISO to provide a framework on which to build a suite of open systems protocols.

Ultimately, the speed at which the TCP/IP-based Internet was adopted, and the rate at which it expanded, caused the development and acceptance of the OSI protocol suite to lag behind.

TCP/IP Model



- Alternative model.
- The architecture of the TCP/IP protocol suite follows the structure of this model.
- Similar to OSI Model

Topic 8: Packet Tracer – Investigating the TCP/IP and OSI Models in Action

This simulation activity is intended to provide a foundation for understanding the TCP/IP protocol suite and the relationship to the OSI model. Simulation mode allows you to view the data contents being sent across the network at each layer.

As data moves through the network, it is broken down into smaller pieces and identified so that the pieces can be put back together when they arrive at the destination. Each piece is assigned a specific name (protocol data unit [PDU]) and associated with a specific layer of the TCP/IP and OSI models. Packet Tracer simulation mode enables you to view each of the layers and the associated PDU. The

following steps lead the user through the process of requesting a web page from a web server by using the web browser application available on a client PC.

Topic 9: Internet Operating System (IOS):

Cisco IOS

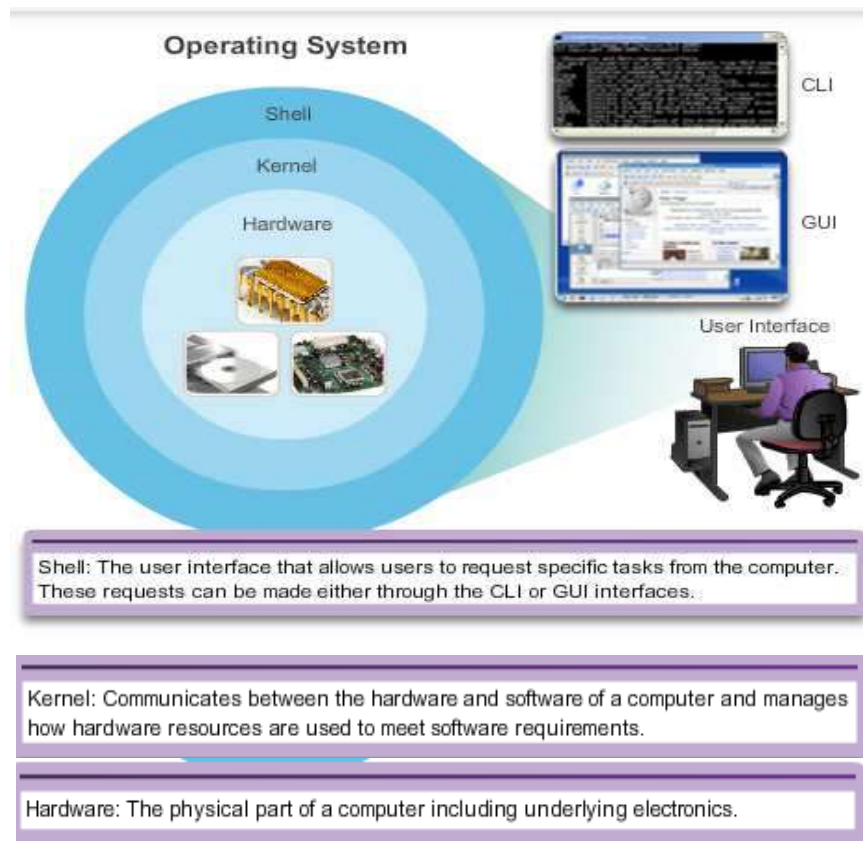
- All networking equipment depend on operating systems:

- End users
- Switches
- Routers
- Wireless access points
- Firewalls

Cisco Internetwork Operating System (IOS)

- Collection of network operating systems used on Cisco devices

Operating System



All end devices and network devices connected to the Internet require an operating system (OS) to help them perform their function.

When a computer is powered on, it loads the OS, normally from a disk drive, into RAM. The portion of the OS code that interacts directly with the computer hardware is known as the kernel. The portion that interfaces with the applications and user is known as the shell. The user can interact with the shell using either the command-line interface (CLI) or graphical user interface (GUI).

When using the CLI, the user interacts directly with the system in a text-based environment by entering commands on the keyboard at a command prompt. The system executes the command, often providing textual output. The GUI interface allows the user to interact with the system in an environment that uses graphical images, multimedia, and text. Actions are performed by interacting with the images on screen. GUI is more user friendly and requires less knowledge of the command structure to utilize the system. For this reason, many individuals rely on the GUI environments. Many operating systems offer both GUI and CLI.

The operating system on home routers is usually called firmware. The most common method for configuring a home router is using a web browser to access an easy to use GUI. Most home routers enable the update of the firmware as new features or security vulnerabilities are discovered.

Infrastructure network devices use a network operating system. The network operating system used on Cisco devices is called the Cisco Internetwork Operating System (IOS). Cisco IOS is a generic term for the collection of network operating systems used on Cisco networking devices. Cisco IOS is used for most Cisco devices regardless of the type or size of the device. The most common method of accessing these devices is using a CLI.

IOS Functions



Cisco IOS routers and switches perform functions that network professionals depend upon to make their networks operate as expected. Major functions performed or enabled by Cisco routers and switches include:

- Providing network security
- IP addressing of virtual and physical interfaces
- Enabling interface-specific configurations to optimize connectivity of the respective media
- Routing

- Enabling quality of service (QoS) technologies
- Supporting network management technologies
- Each feature or service has an associated collection of configuration commands that allow a network technician to implement it.
- The services provided by the Cisco IOS are generally accessed using a CLI.

Topic 10: Accessing an IOS Device:

There are several ways to access the CLI environment. The most common methods are:

Console



The console port is a management port that provides out-of-band access to Cisco device. Out-of-band access refers to access via a dedicated management channel that is used for device maintenance purposes only. The advantage of using a console port is that the device is accessible even if no networking services have been configured, such as when performing an initial configuration of the networking device. When performing an initial configuration, a computer running terminal emulation software is connected to the console port of the device using a special cable. Configuration commands for setting up the switch or router can be entered on the connected computer.

The console port can also be used when the networking services have failed and remote access of the Cisco IOS device is not possible. If this occurs, a connection to the console can enable a computer to determine the status of the device. By default, the console conveys the device startup, debugging, and error messages. After the network technician is connected to the device, the network technician can perform any configuration commands necessary using the console session.

For many IOS devices, console access does not require any form of security, by default. However, the console should be configured with passwords to prevent unauthorized device access. In the event that a password is lost, there is a special set of procedures for bypassing the password and accessing the device. The device should also be located in a locked room or equipment rack to prevent unauthorized physical access.

Telnet

Telnet is a method for remotely establishing a CLI session of a device, through a virtual interface, over a network. Unlike the console connection, Telnet sessions require active networking services on the device. The network device must have at least one active interface configured with an Internet address, such as an IPv4 address. Cisco IOS devices include a Telnet server process that allows users to enter configuration commands from a Telnet client. In addition to supporting the Telnet server process, the Cisco IOS device also contains a Telnet client. This allows a network administrator to telnet from the Cisco device CLI to any other device that supports a Telnet server process.

SSH

The Secure Shell (SSH) protocol provides a remote login similar to Telnet, except that it uses more secure network services. SSH provides stronger password authentication than Telnet and uses encryption when transporting session data. This keeps the user ID, password, and the details of the management session private. As a best practice, use SSH instead of Telnet whenever possible.

Most versions of Cisco IOS include an SSH server. In some devices, this service is enabled by default. Other devices require the SSH server to be enabled manually. IOS devices also include an SSH client that can be used to establish SSH sessions with other devices.

AUX

An older way to establish a CLI session remotely is via a telephone dialup connection using a modem connected to the auxiliary (AUX) port of a router. Similar to the console connection, the AUX method is also an out-of-band connection and does not require any networking services to be configured or available on the device. In the event that network services have failed, it may be possible for a remote administrator to access the **switch** or router over a telephone line.

The AUX port can also be used locally, like the console port, with a direct connection to a computer running a terminal emulation program. However, the console port is preferred over the AUX port for troubleshooting because it displays startup, debugging, and error messages by default.

Terminal Emulation Program

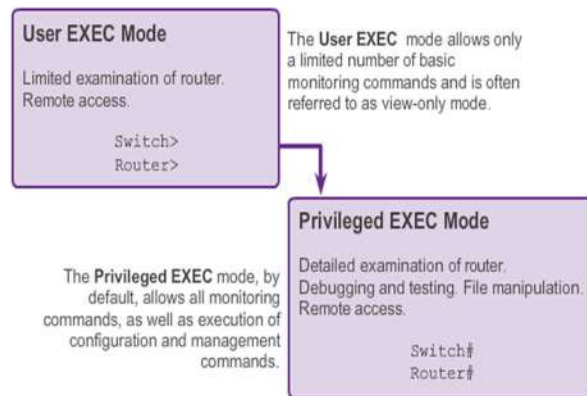
Software available for connecting to a networking device (usually same as terminal/serial/console connection):

- PuTTY
- Tera Term
- HyperTerminal
- OS X Terminal

Topic 11: IOS Modes of Operation:

After a network technician is connected to a device, it is possible to configure it. The network technician must navigate through various modes of the IOS. The Cisco IOS modes are quite similar for switches and routers. The CLI uses a hierarchical structure for the modes.

In hierarchical order from most basic to most specialized, the major modes are:



The two primary modes of operation are user EXEC mode and privileged EXEC mode. As a security feature, the Cisco IOS software separates the EXEC sessions into two levels of access. The privileged EXEC mode has a higher level of authority in what it allows the user to do with the device.

User EXEC Mode

The user EXEC mode has limited capabilities but is useful for some basic operations. The user EXEC mode is at the most basic level of the modal hierarchical structure. This mode is the first mode encountered upon entrance into the CLI of an IOS device.

The user EXEC mode allows only a limited number of basic monitoring commands. This is often referred to as view-only mode. The user EXEC level does not allow the execution of any commands that might change the configuration of the device.

By default, there is no authentication required to access the user EXEC mode from the console. However, it is a good practice to ensure that authentication is configured during the initial configuration.

The user EXEC mode is identified by the CLI prompt that ends with the > symbol. This is an example that shows the > symbol in the prompt:

```
Switch>
```

Privileged EXEC Mode

The execution of configuration and management commands requires that the network administrator use the privileged EXEC mode or a more specific mode in the hierarchy. This means that a user must enter user EXEC mode first, and from there, access privileged EXEC mode.

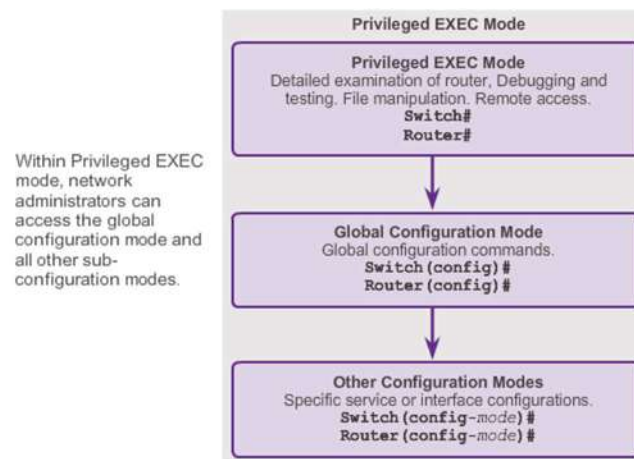
The privileged EXEC mode can be identified by the prompt ending with the # symbol.

```
Switch#
```

By default, privileged EXEC mode does not require authentication. It is a good practice to ensure that authentication is configured.

Global configuration mode and all other more specific configuration modes can only be reached from the privileged EXEC mode. In a later section of this chapter, we will examine device configuration and some of the configuration modes.

Global configuration mode



Global configuration mode and interface configuration modes can only be reached from the privileged EXEC mode.

Global Configuration Mode

The primary configuration mode is called global configuration or global config. From global configuration mode, CLI configuration changes are made that affect the operation of the device as a whole. The global configuration mode is accessed before accessing specific configuration modes.

The following CLI command is used to take the device from privileged EXEC mode to the global configuration mode and to allow entry of configuration commands from a terminal:

Switch# **configure terminal**

After the command is executed, the prompt changes to show that the switch is in global configuration mode.

Switch(config)#

Specific Configuration Modes

From the global configuration mode, the user can enter different sub-configuration modes. Each of these modes allows the configuration of a particular part or function of the IOS device. The list below shows a few of them:

Interface mode - to configure one of the network interfaces (Fa0/0, S0/0/0)

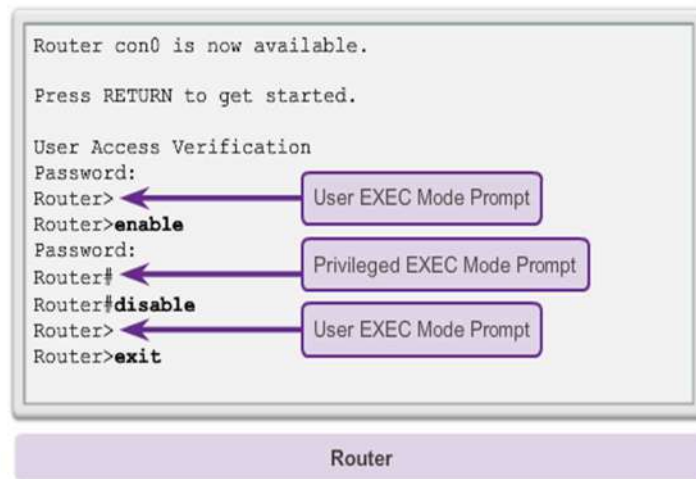
Line mode - to configure one of the physical or virtual lines (console, AUX, VTY)

Command Prompts

When using the CLI, the mode is identified by the command-line prompt that is unique to that mode. By default, every prompt begins with the device name. Following the name, the remainder of the prompt indicates the mode. For example, the default prompt for the global configuration mode on a switch would be:

Switch(config)#

Navigating between IOS Modes



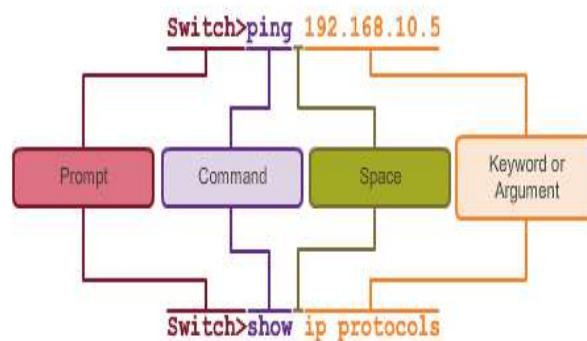
Moving Between the User EXEC and Privileged EXEC Modes

The **enable** and **disable** commands are used to change the CLI between the user EXEC mode and the privileged EXEC mode, respectively.

In order to access the privileged EXEC mode, use the **enable** command. The privileged EXEC mode is sometimes called the enable mode.

Topic 12: The Command Structure:

Basic IOS Command Structure

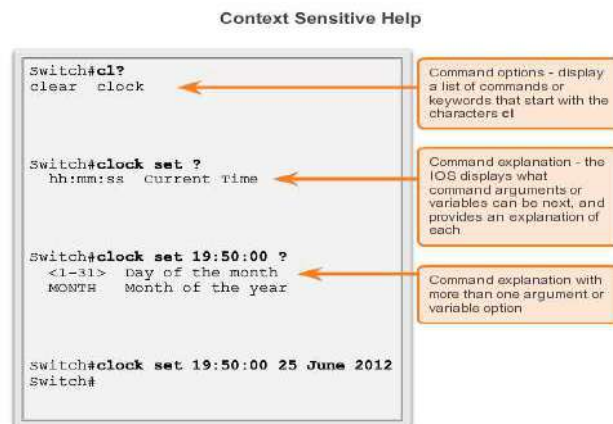


A Cisco IOS device supports many commands. Each IOS command has a specific format or syntax and can only be executed at the appropriate mode. The general syntax for a command is the command followed by any appropriate keywords and arguments. Some commands include a subset of keywords and arguments that provide additional functionality. Commands are used to execute an action, and the keywords are used to identify where or how to execute the command.

As shown in Figure above, the command is the initial word or words entered in the command line following the prompt. The commands are not case-sensitive. Following the command are one or more

keywords and arguments. After entering each complete command, including any keywords and arguments, press the Enter key to submit the command to the command interpreter.

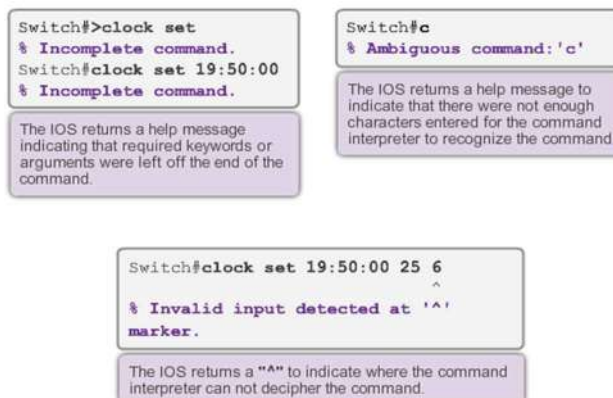
Context-Sensitive Help



The context-sensitive help provides a list of commands and the arguments associated with those commands within the context of the current mode. To access context-sensitive help, enter a question mark, `?`, at any prompt. There is an immediate response without the need to use the Enter key.

One use of context-sensitive help is to get a list of available commands. This can be used when you are unsure of the name for a command or you want to see if the IOS supports a particular command in a particular mode.

Command Syntax Check



When a command is submitted by pressing the Enter key, the command line interpreter parses the command from left to right to determine what action is being requested. The IOS generally only provides negative feedback, as shown in Figure above. If the interpreter understands the command, the requested action is executed and the CLI returns to the appropriate prompt. However, if the interpreter cannot understand the command being entered, it will provide feedback describing what is wrong with the command.

There are three different types of error messages:

Ambiguous command

Incomplete command

Incorrect command

Hot Keys and Shortcuts

The IOS CLI provides hot keys and shortcuts that make configuring, monitoring, and troubleshooting easier.

The following are worthy of special note:

Down Arrow - Allows the user to scroll forward through former commands

Up Arrow - Allows the user to scroll backward through former commands

Tab - Completes the remainder of a partially typed command or keyword

Ctrl-A - Moves to the beginning of the line

Ctrl-E - Moves to the end of the line

Ctrl-R - Redisplays a line

Ctrl-Z - Exits the configuration mode and returns to user EXEC

Ctrl-C - Exits the configuration mode or aborts the current command

Ctrl-Shift-6 - Allows the user to interrupt an IOS process such as ping or traceroute

Abbreviated commands or keywords

Commands and keywords can be abbreviated to the minimum number of characters that identify a unique selection. For example, the **configure** command can be abbreviated to **conf** because **configure** is the only command that begins with **conf**. An abbreviation of **con** will not work because more than one command begins with **con**.

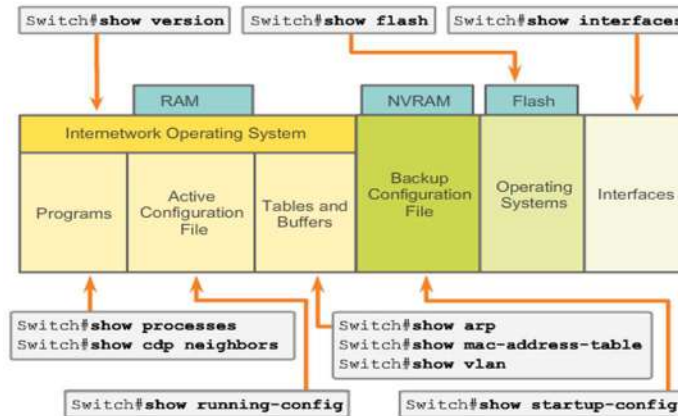
Keywords can also be abbreviated.

As another example, **show interfaces** can be abbreviated like this:

Switch# **show interfaces**

Switch# **show int**

IOS Examination Commands



One of the most commonly used commands on a switch or router is:

Switch# **show version**

This command displays information about the currently loaded IOS version, along with hardware and device information. If you are logged into a router or switch remotely, the **show version** command is an excellent means of quickly finding useful summary information about the particular device to which you are connected.

Topic 13: Packet Tracer - Navigating the IOS:

In this activity, you will practice skills necessary for navigating the Cisco IOS, including different user access modes, various configuration modes, and common commands you use on a regular basis.

Topic 14: Configuring Hostnames:

Cisco switches and Cisco routers have many similarities. They support a similar modal operating system support similar command structures, and support many of the same commands. In addition, both devices have identical initial configuration steps when implementing them in a network.

However, a Cisco IOS switch is one of the simplest devices that can be configured on a network. This is because there are no configurations that are required prior to the device functioning. At its most basic, a switch can be plugged in with no configuration, but it will still switch data between connected devices.

A switch is also one of the fundamental devices used in the creation of a small network. By connecting two PCs to a switch, those PCs will instantly have connectivity with one another. Initial settings include setting a name for the switch, limiting access to the device configuration, configuring banner messages, and saving the configuration.

Device Names

When configuring a networking device, one of the first steps is configuring a unique device name, or hostname. Hostnames appear in CLI prompts, can be used in various authentication processes between devices, and should be used on topology diagrams.

Hostnames are configured on the active networking device. If the device name is not explicitly configured, a factory-assigned default device name is used by Cisco IOS. The default name for a Cisco IOS switch is "Switch."

Some guidelines for naming conventions are that names should:

- Start with a letter
- Contain no spaces
- End with a letter or digit
- Use only letters, digits, and dashes
- Be less than 64 characters in length

Hostnames allow devices to be identified by network administrators over a network or the Internet.

Securing Privilege EXEC Mode

```
Sw-Floor-1>enable
Sw-Floor-1#
Sw-Floor-1#conf terminal
Sw-Floor-1(config)#enable secret class
Sw-Floor-1(config)#exit
Sw-Floor-1#
Sw-Floor-1#disable
Sw-Floor-1>enable
Password:
Sw-Floor-1#
```

- use the **enable secret** command, not the older **enable password** command
- **enable secret** provides greater security because the password is encrypted

Securing USER EXEC Mode

```
Sw-Floor-1(config)#line console 0
Sw-Floor-1(config-line)#password cisco
Sw-Floor-1(config-line)#login
Sw-Floor-1(config-line)#exit
Sw-Floor-1(config)#
Sw-Floor-1(config)#line vty 0 15
Sw-Floor-1(config-line)#password cisco
Sw-Floor-1(config-line)#login
Sw-Floor-1(config-line)#
```

- **Console port** must be secured
 - Reduces the chance of unauthorized personnel physically plugging a cable into the device and gaining device access
- **VTY lines** allow access to a Cisco device via Telnet

Topic 15: Packet Tracer - Configuring Initial Switch:

In this activity, you will perform basic switch configurations. You will secure access to the command-line interface (CLI) and console ports using encrypted and plain text passwords. You will also learn how to configure messages for users logging into the switch. These banners are also used to warn unauthorized users that access is prohibited.

Topic 16: Packet Tracer - Implementing Basic Connectivity:

In this activity, you will first perform basic switch configurations. Then you will implement basic connectivity by configuring IP addressing on switches and PCs. When the IP addressing configuration is complete, you will use various **show** commands to verify configurations and use the **ping** command to verify basic connectivity between devices.

Topic 17: Packet Tracer - Implementing Basic Connectivity - 2:

This activity is in continuation of previous topic.

Topic 18: Packet Tracer - Configuring Switch Management Address:

In this activity, you will first perform basic switch configurations. Then you will implement basic connectivity by configuring IP addressing on switches and PCs. When the IP addressing configuration is complete, you will use various **show** commands to verify configurations and use the **ping** command to verify basic connectivity between devices.

Topic 19: Physical Layer Protocols:

Whether connecting to a local printer in the home or to a web site in another country, before any network communications can occur, a physical connection to a local network must be established first. A physical connection can be a wired connection using a cable or a wireless connection using radio waves.

The type of physical connection used is totally dependent upon the setup of the network.

Network Interface Cards (NICs)

Network Interface Cards (NICs) connect a device to the network. Ethernet NICs are used for a wired connection whereas WLAN (Wireless Local Area Network) NICs are used for wireless. An end-user device may include one or both types of NICs. A network printer, for example, may only have an Ethernet NIC, and therefore, must connect to the network using an Ethernet cable. Other devices, such as tablets and smart phones, might only contain a WLAN NIC and must use a wireless connection.

There are three basic forms of network media. The physical layer produces the representation and groupings of bits for each type of media as:

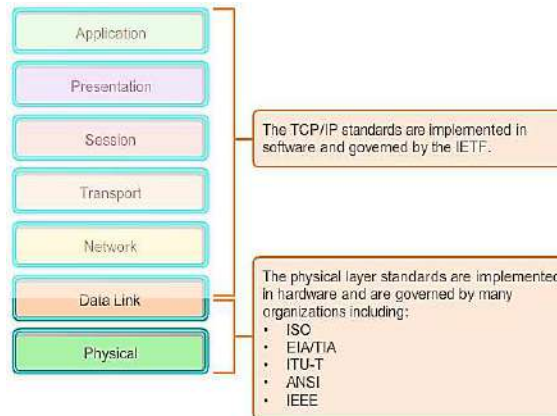
Copper cable: The signals are patterns of electrical pulses.

Fiber-optic cable: The signals are patterns of light.

Wireless: The signals are patterns of microwave transmissions.

To enable physical layer interoperability, all aspects of these functions are governed by standard organizations.

Physical Layer Standards



Bandwidth

Different physical media support the transfer of bits at different speeds. Data transfer is usually discussed in terms of bandwidth and throughput.

Bandwidth is the capacity of a medium to carry data. Digital bandwidth measures the amount of data that can flow from one place to another in a given amount of time. Bandwidth is typically measured in kilobits per second (kb/s) or megabits per second (Mb/s).

The practical bandwidth of a network is determined by a combination of factors:

- The properties of the physical media
- The technologies chosen for signaling and detecting network signals
- Physical media properties, current technologies, and the laws of physics all play a role in determining available bandwidth.

Throughput

Throughput is the measure of the transfer of bits across the media over a given period of time.

Due to a number of factors, throughput usually does not match the specified bandwidth in physical layer implementations. Many factors influence throughput including:

- The amount of traffic
- The type of traffic
- The latency created by the number of network devices encountered between source and destination

Latency refers to the amount of time, to include delays, for data to travel from one given point to another.

In an internetwork or network with multiple segments, throughput cannot be faster than the slowest link of the path from source to destination. Even if all or most of the segments have high bandwidth, it will only take one segment in the path with low throughput to create a bottleneck to the throughput of the entire network.

Note: There is a third measurement to measure the transfer of usable data that is known as goodput. Goodput is the measure of usable data transferred over a given period of time. Goodput is throughput minus traffic overhead for establishing sessions, acknowledgements, and encapsulation.

Topic 20: Network Media:

There are three main types of copper media used in networking:

Unshielded Twisted-Pair (UTP)

Unshielded twisted-pair (UTP) cabling is the most common networking media. UTP cabling, terminated with RJ-45 connectors, is used for interconnecting network hosts with intermediate networking devices, such as switches and routers.

In LANs, UTP cable consists of four pairs of color-coded wires that have been twisted together and then encased in a flexible plastic sheath which protects from minor physical damage. The twisting of wires helps protect against signal interference from other wires.

Shielded Twisted-Pair (STP)

Shielded twisted-pair (STP) provides better noise protection than UTP cabling. However, compared to UTP cable, STP cable is significantly more expensive and difficult to install. Like UTP cable, STP uses an RJ-45 connector.

STP cable combines the techniques of shielding to counter EMI and RFI and wire twisting to counter crosstalk. To gain the full benefit of the shielding, STP cables are terminated with special shielded STP data connectors. If the cable is improperly grounded, the shield may act like an antenna and pick up unwanted signals.

Coaxial

Coaxial cable, or coax for short, gets its name from the fact that there are two conductors that share the same axis. Coaxial cable consists of:

- A copper conductor used to transmit the electronic signals.
- The copper conductor is surrounded by a layer of flexible plastic insulation.
- The insulating material is surrounded in a woven copper braid, or metallic foil, that acts as the second wire in the circuit and as a shield for the inner conductor. This second layer, or shield, also reduces the amount of outside electromagnetic interference.
- The entire cable is covered with a cable jacket to protect it from minor physical damage.

Fiber Optic Cabling

Optical fiber cable has become very popular for interconnecting infrastructure network devices. It permits the transmission of data over longer distances and at higher bandwidths (data rates) than any other networking media.

Optical fiber is a flexible but extremely thin transparent strand of very pure glass (silica) not much bigger than a human hair. Bits are encoded on the fiber as light impulses. The fiber-optic cable acts as a waveguide, or light pipe to transmit light between the two ends with minimal loss of signal.

Unlike copper wires, fiber-optic cable can transmit signals with less attenuation and is completely immune to EMI and RFI.

Fiber-optic cabling is now being used in four types of industry:

Enterprise Networks: Fiber is used for backbone cabling applications and interconnecting infrastructure devices.

FTTH and Access Networks: Fiber-to-the-home (FTTH) is used to provide always-on broadband services to homes and small businesses. FTTH supports affordable high-speed Internet access, as well as telecommuting, telemedicine, and video on demand.

Long-Haul Networks: Service providers use long-haul terrestrial optical fiber networks to connect countries and cities. Networks typically range from a few dozen to a few thousand kilometers and use up to 10 Gb/s-based systems.

Submarine Networks: Special fiber cables are used to provide reliable high-speed, high-capacity solutions capable of or surviving in harsh undersea environments up to transoceanic distances.

Our focus is the use of fiber within the enterprise.

Wireless Media

	• IEEE 802.11 standards • Commonly referred to as Wi-Fi • Uses CSMA/CA • Variations include: • 802.11a: 54 Mb/s, 5 GHz • 802.11b: 11 Mb/s, 2.4 GHz • 802.11g: 54 Mb/s, 2.4 GHz • 802.11n: 600 Mb/s, 2.4, and 5 GHz • 802.11ac: 1 Gb/s, 5 GHz • 802.11ad: 7 Gb/s, 2.4 GHz, 5 GHz, and 60 GHz
	• IEEE 802.15 standard • Supports speeds up to 3 Mb/s • Provides device pairing over distances from 1 to 100 meters
	• IEEE 802.16 standard • Provides speeds up to 1 Gb/s • Uses a point-to-multipoint topology to provide wireless broadband access

The IEEE and telecommunications industry standards for wireless data communications cover both the data link and physical layers.

Topic 21: Data Link Layer Protocols:

The TCP/IP network access layer is the equivalent of the OSI:

Data link (Layer 2)

Physical (Layer 1)

The data link layer is responsible for the exchange of frames between nodes over a physical network media. It allows the upper layers to access the media and controls how data is placed and received on the media.

Specifically the data link layer performs these two basic services:

- It accepts Layer 3 packets and packages them into data units called frames.
- It controls media access control and performs error detection.

The data link layer effectively separates the media transitions that occur as the packet is forwarded from the communication processes of the higher layers. The data link layer receives packets from and directs packets to an upper layer protocol, in this case IPv4 or IPv6. This upper layer protocol does not need to be aware of which media the communication will use.

Data Link Sublayers

The data link layer is actually divided into two sublayers:

Logical Link Control (LLC): This upper sublayer defines the software processes that provide services to the network layer protocols. It places information in the frame that identifies which network layer protocol is being used for the frame. This information allows multiple Layer 3 protocols, such as IPv4 and IPv6, to utilize the same network interface and media.

Media Access Control (MAC): This lower sublayer defines the media access processes performed by the hardware. It provides data link layer addressing and delimiting of data according to the physical signaling requirements of the medium and the type of data link layer protocol in use.

Separating the data link layer into sublayers allows for one type of frame defined by the upper layer to access different types of media defined by the lower layer. Such is the case in many LAN technologies, including Ethernet.

Data Link Frame Fields

The frame header contains the control information specified by the data link layer protocol for the specific logical topology and media used.

Frame control information is unique to each type of protocol. It is used by the Layer 2 protocol to provide features demanded by the communication environment.

Start Frame field: Indicates the beginning of the frame.

Source and Destination Address fields: Indicates the source and destination nodes on the media.

Type field: Indicates the upper layer service contained in the frame.

Different data link layer protocols may use different fields from those mentioned. For example other Layer 2 protocol header frame fields could include:

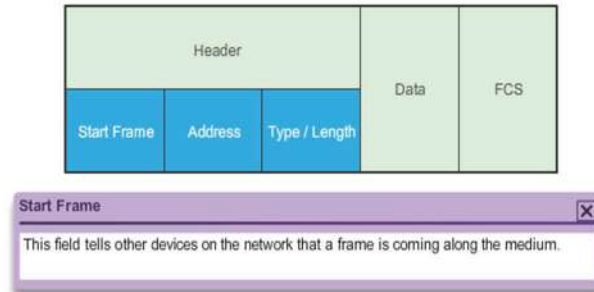
Priority/Quality of Service field: Indicates a particular type of communication service for processing.

Logical connection control field: Used to establish a logical connection between nodes.

Physical link control field: Used to establish the media link.

Flow control field: Used to start and stop traffic over the media.

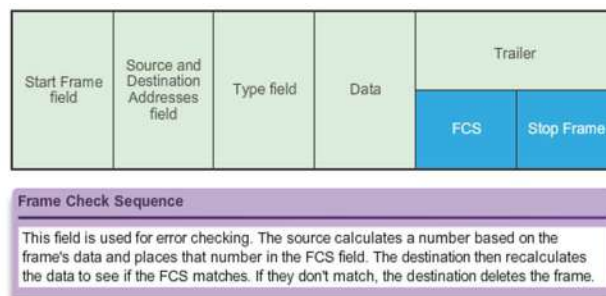
Congestion control field: Indicates congestion in the media.



Data link layer protocols add a trailer to the end of each frame. The trailer is used to determine if the frame arrived without error. This process is called error detection and is accomplished by placing a logical or mathematical summary of the bits that comprise the frame in the trailer. Error detection is added at the data link layer because the signals on the media could be subject to interference, distortion, or loss that would substantially change the bit values that those signals represent.

A transmitting node creates a logical summary of the contents of the frame. This is known as the cyclic redundancy check (CRC) value. This value is placed in the Frame Check Sequence (FCS) field of the frame to represent the contents of the frame.

When the frame arrives at the destination node, the receiving node calculates its own logical summary, or CRC, of the frame. The receiving node compares the two CRC values. If the two values are the same, the frame is considered to have arrived as transmitted. If the CRC value in the FCS differs from the CRC calculated at the receiving node, the frame is discarded.



Therefore, the FCS field is used to determine if errors occurred in the transmission and reception of the frame. The error detection mechanism provided by the use of the FCS field discovers most errors caused on the media.

Topic 22: Packet Tracer – Connecting a Wired and Wireless LAN:

When working in Packet Tracer (a lab environment or a corporate setting), you should know how to select the appropriate cable and how to properly connect devices. This activity will examine device configurations in Packet Tracer, selecting the proper cable based on the configuration, and connecting the devices. This activity will also explore the physical view of the network in Packet Tracer.

Topic 23: Network Layer Protocols:

The network layer, or OSI Layer 3, provides services to allow end devices to exchange data across the network. To accomplish this end-to-end transport, the network layer uses four basic processes:

Addressing end devices - In the same way that a phone has a unique telephone number, end devices must be configured with a unique IP address for identification on the network. An end device with a configured IP address is referred to as a host.

Encapsulation - The network layer receives a protocol data unit (PDU) from the transport layer. In a process called encapsulation, the network layer adds IP header information, such as the IP address of the source (sending) and destination (receiving) hosts. After header information is added to the PDU, the PDU is called a packet.

Routing - The network layer provides services to direct packets to a destination host on another network. To travel to other networks, the packet must be processed by a router. The role of the router is to select paths for and direct packets toward the destination host in a process known as routing. A packet may cross many intermediary devices before reaching the destination host. Each route the packet takes to reach the destination host is called a hop.

De-encapsulation - When the packet arrives at the network layer of the destination host, the host checks the IP header of the packet. If the destination IP address within the header matches its own IP address, the IP header is removed from the packet. This process of removing headers from lower layers is known as de-encapsulation. After the packet is de-encapsulated by the network layer, the resulting Layer 4 PDU is passed up to the appropriate service at the transport layer.

Unlike the transport layer (OSI Layer 4), which manages the data transport between the processes running on each host, network layer protocols specify the packet structure and processing used to carry the data from one host to another host. Operating without regard to the data carried in each packet allows the network layer to carry packets for multiple types of communications between multiple hosts.

There are several network layer protocols in existence; however, only the following two are commonly implemented:

- Internet Protocol version 4 (IPv4)
- Internet Protocol version 6 (IPv6)

Other legacy network layer protocols that are not widely used include:

- Novell Internetwork Packet Exchange (IPX)
- AppleTalk
- Connectionless Network Service (CLNS/DECNet)

Characteristics of IP Protocol

IP is the network layer service implemented by the TCP/IP protocol suite.

IP was designed as a protocol with low overhead. It provides only the functions that are necessary to deliver a packet from a source to a destination over an interconnected system of networks. The protocol was not designed to track and manage the flow of packets. These functions, if required, are performed by other protocols in other layers.

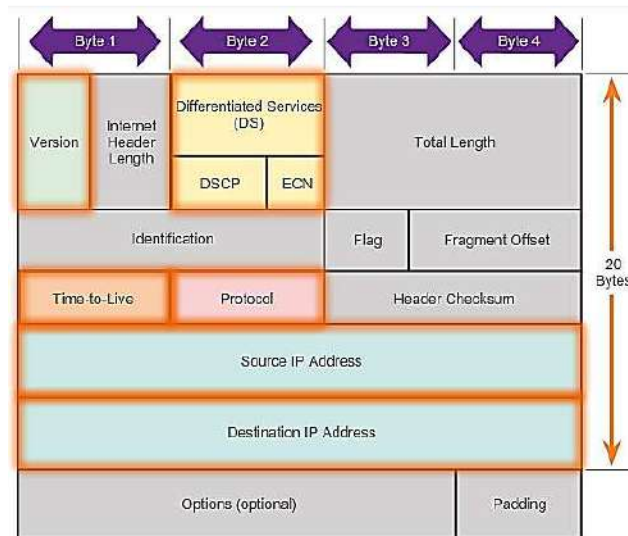
The basic characteristics of IP are:

Connectionless - No connection with the destination is established before sending data packets.

Best Effort (unreliable) - Packet delivery is not guaranteed.

Media Independent - Operation is independent of the medium carrying the data.

Topic 24: IPv4 Packet:



An IPv4 packet has two parts:

IP Header - Identifies the packet characteristics.

Payload - Contains the Layer 4 segment information and the actual data.

As shown in the figure above, an IPv4 packet header consists of fields containing important information about the packet. These fields contain binary numbers which are examined by the Layer 3 process. The binary values of each field identify various settings of the IP packet.

Significant fields in the IPv4 header include:

Version - Contains a 4-bit binary value identifying the IP packet version. For IPv4 packets, this field is always set to 0100.

Differentiated Services (DS) -Formerly called the Type of Service (ToS) field, the DS field is an 8-bit field used to determine the priority of each packet. The first 6 bits identify the Differentiated Services Code Point (DSCP) value that is used by a quality of service (QoS) mechanism. The last 2 bits identify the explicit congestion notification (ECN) value that can be used to prevent dropped packets during times of network congestion.

Time-to-Live (TTL) - Contains an 8-bit binary value that is used to limit the lifetime of a packet. It is specified in seconds but is commonly referred to as hop count. The packet sender sets the initial time-to-live (TTL) value and is decreased by one each time the packet is processed by a router, or hop. If the TTL field decrements to zero, the router discards the packet and sends an Internet Control Message Protocol (ICMP) Time Exceeded message to the source IP address. The **traceroute** command uses this field to identify the routers used between the source and destination.

Protocol - This 8-bit binary value indicates the data payload type that the packet is carrying, which enables the network layer to pass the data to the appropriate upper-layer protocol. Common values include ICMP (0x01), TCP (0x06), and UDP (0x11).

Source IP Address - Contains a 32-bit binary value that represents the source IP address of the packet.

Destination IP Address - Contains a 32-bit binary value that represents the destination IP address of the packet.

Topic 25: IPv4 Address and Subnet Mask:

Understanding binary notation is important when determining if two hosts are in the same network. Recall that an IP address is a hierarchical address that is made up of two parts: a network portion and a host portion. But when determining the network portion versus the host portion, it is necessary to look, not at the decimal value, but at the 32-bit stream. Within the 32-bit stream, a portion of the bits makes up the network and a portion of the bits makes up the host.

The bits within the network portion of the address must be identical for all devices that reside in the same network. The bits within the host portion of the address must be unique to identify a specific host within a network. Regardless of whether the decimal numbers between two IPv4 addresses match up, if two hosts have the same bit-pattern in the specified network portion of the 32-bit stream, those two hosts will reside in the same network.

But how do hosts know which portion of the 32-bits is network and which is host? That is the job of the subnet mask.

When an IP host is configured, a subnet mask is assigned along with an IP address. Like the IP address, the subnet mask is 32 bits long. The subnet mask signifies which part of the IP address is network and which part is host.

The subnet mask is compared to the IP address from left to right, bit for bit. The 1s in the subnet mask represent the network portion; the 0s represent the host portion.

Similar to IPv4 addresses, the subnet mask is represented in dotted decimal format for ease of use. The subnet mask is configured on a host device, in conjunction with the IPv4 address, and is required so the host can determine which network it belongs to.

Network Prefixes

The prefix length is another way of expressing the subnet mask. The prefix length is the number of bits set to 1 in the subnet mask. It is written in slash notation, a / followed by the number of bits set to 1. For example, if the subnet mask is 255.255.255.0, there are 24 bits set to 1 in the binary version of the subnet mask, so the prefix length is 24 bits or /24. The prefix and the subnet mask are different ways of representing the same thing - the network portion of an address.

IPv4 Subnet Mask

When an IPv4 address is assigned to a device, that device uses the subnet mask to determine what network address the device belongs to. The network address is the address that represents all the devices on the same network.

When sending network data, the device uses this information to determine whether it can send packets locally, or if it must send the packets to a default gateway for remote delivery. When a host sends a packet, it compares the network portion of its own IP address to the network portion of the destination IP address, based on subnet masks. If the network bits match, both the source and destination host are on the same network and the packet can be delivered locally. If they do not match, the sending host forwards the packet to the default gateway to be sent on to the other network.

When an IPv4 address is assigned to a device, that device uses the subnet mask to determine what network address the device belongs to. The network address is the address that represents all the devices on the same network.

When sending network data, the device uses this information to determine whether it can send packets locally, or if it must send the packets to a default gateway for remote delivery. When a host sends a packet, it compares the network portion of its own IP address to the network portion of the destination IP address, based on subnet masks. If the network bits match, both the source and destination host are on the same network and the packet can be delivered locally. If they do not match, the sending host forwards the packet to the default gateway to be sent on to the other network.

There are three types of addresses within the address range of each IPv4 network:

- Network address
- Host addresses
- Broadcast address

Network Address

The network address is a standard way to refer to a network. The subnet mask or the prefix length might also be used when referring to network address.

Host Address

Every end device requires a unique address to communicate on the network. In IPv4 addresses, the values between the network address and the broadcast address can be assigned to end devices in a network.

Broadcast Address

The IPv4 broadcast address is a special address for each network that allows communication to all the hosts in that network. To send data to all hosts in a network at once, a host can send a single packet that is addressed to the broadcast address of the network, and each host in the network that receives this packet will process its contents.

First Host Address

The host portion of the first host address will contain all 0 bits with a 1 bit for the lowest order or right-most bit. This address is always one greater than the network address. In this example the first host address on the 10.1.1.0/24 network is 10.1.1.1. It is common in many addressing schemes to use the first host address for the router or default gateway address.

Last Host Address

The host portion of the last host address will contain all 1 bits with a 0 bit for the lowest order or right-most bit. This address is always one less than the broadcast address. The last host address on the 10.1.1.0/24 network is 10.1.1.254.

Topic 26: IPv4 Unicast, Broadcast and Multicast:

In an IPv4 network, the hosts can communicate one of three ways:

Unicast - The process of sending a packet from one host to an individual host

Broadcast - The process of sending a packet from one host to all hosts in the network

Multicast - The process of sending a packet from one host to a selected group of hosts, possibly in different networks

These three types of communication are used for different purposes in data networks. In all three cases, the IPv4 address of the originating host is placed in the packet header as the source address.

Unicast Traffic

In an IPv4 network, the unicast addresses applied to an end device is referred to as the host address. For unicast communication, the addresses assigned to the two end devices are used as the source and destination IPv4 addresses. During the encapsulation process, the source host places its IPv4 address in the unicast packet header as the source address and the IPv4 address of the destination host in the packet header as the destination address. Regardless of whether the destination specified a packet is a unicast, broadcast or multicast; the source address of any packet is always the unicast address of the originating host.

IPv4 host addresses are unicast addresses and are in the address range of 0.0.0.0 to 223.255.255.255. However, within this range are many addresses that are reserved for special purposes. These special purpose addresses will be discussed later in this chapter.

Multicast Transmission

Multicast transmission is designed to conserve the bandwidth of an IPv4 network. It reduces traffic by allowing a host to send a single packet to a selected set of hosts that are part of a subscribing multicast group.

Some examples of multicast transmission are:

- Video and audio broadcasts
- Routing information exchange by routing protocols
- Distribution of software
- Remote gaming

Multicast Addresses

IPv4 has a block of addresses reserved for addressing multicast groups. This address range is 224.0.0.0 to 239.255.255.255. The multicast address range is subdivided into different types of addresses:

reserved link local addresses and globally scoped addresses. One additional type of multicast address is the administratively scoped addresses, also called limited scope addresses.

Multicast Clients

Hosts that receive particular multicast data are called multicast clients. The multicast clients use services requested by a client program to subscribe to the multicast group.

Each multicast group is represented by a single IPv4 multicast destination address. When an IPv4 host subscribes to a multicast group, the host processes packets addressed to this multicast address and packets addressed to its uniquely allocated unicast address.

Broadcast Transmission

Broadcast traffic is used to send packets to all hosts in the network using the broadcast address for the network. With a broadcast, the packet contains a destination IP address with all ones (1s) in the host portion. This means that all hosts on that local network (broadcast domain) will receive and look at the packet. Many network protocols, such as DHCP, use broadcasts. When a host receives a packet sent to the network broadcast address, the host processes the packet as it would a packet addressed to its unicast address.

Some examples for using broadcast transmission are:

- Mapping upper layer addresses to lower layer addresses
- Requesting an address

Unlike unicast, where the packets can be routed throughout the internetwork, broadcast packets are usually restricted to the local network. This restriction is dependent on the configuration of the gateway router and the type of broadcast. There are two types of broadcasts: directed broadcast and limited broadcast.

Directed Broadcast

A directed broadcast is sent to all hosts on a specific network. This type of broadcast is useful for sending a broadcast to all hosts on a non-local network. For example, for a host outside of the 172.16.4.0/24 network to communicate with all of the hosts within that network, the destination address of the packet would be 172.16.4.255. Although routers do not forward directed broadcasts by default, they may be configured to do so.

Limited Broadcast

The limited broadcast is used for communication that is limited to the hosts on the local network. These packets always use a destination IPv4 address 255.255.255.255. Routers do not forward a limited broadcast. For this reason, an IPv4 network is also referred to as a broadcast domain. Routers form the boundary for a broadcast domain.

As an example, a host within the 172.16.4.0/24 network would broadcast to all hosts in its network using a packet with a destination address of 255.255.255.255.

Topic 27: Packet Tracer – Investigate Unicast, Broadcast and Multicast Traffic :

This activity will examine unicast, broadcast, and multicast behavior. Most traffic in a network is unicast. When a PC sends an ICMP echo request to a remote router, the source address in the IP packet header is the IP address of the sending PC. The destination address in the IP packet header is the IP address of the interface on the remote router. The packet is sent only to the intended destination.

Using the **ping** command or the Add Complex PDU feature of Packet Tracer, you can directly ping broadcast addresses to view broadcast traffic.

Topic 28: Types of IPv4 Addresses:

Although most IPv4 host addresses are public addresses designated for use in networks that are accessible on the Internet, there are blocks of addresses that are used in networks that require limited or no Internet access. These addresses are called private addresses.

Private Addresses

The private address blocks are:

10.0.0.0 to 10.255.255.255 (10.0.0.0/8)

172.16.0.0 to 172.31.255.255 (172.16.0.0/12)

192.168.0.0 to 192.168.255.255 (192.168.0.0/16)

Public Addresses

Most IPv4 addresses are public IP addresses.

These are reachable IPv4 addresses on the Internet.

However, there are blocks of addresses that are private addresses and are never propagated on the Internet.

Packets with a ***source or destination private IP address are not propagated by Internet routers.***

Internet routers / firewalls block or translate these addresses.

Private addresses are usually converted to public IP addresses using **NAT (Network Address Translation)**

- Private addresses are defined in RFC 1918.
- The private address blocks are:
 - **10.0.0.0 /8:** 10.0.0.0 to 10.255.255.255
 - **172.16.0.0 /12:** 172.16.0.0 to 172.31.255.255
 - **192.168.0.0 /16:** 192.168.0.0 to 192.168.255.255

Special Addresses

There are certain addresses that cannot be assigned to hosts. There are also special addresses that can be assigned to hosts, but with restrictions on how those hosts can interact within the network.

Network and Broadcast Addresses

Loopback

One such reserved address is the IPv4 loopback address 127.0.0.1. The loopback is a special address that hosts use to direct traffic to themselves. The loopback address creates a shortcut method for TCP/IP applications and services that run on the same device to communicate with one another. By using the loopback address instead of the assigned IPv4 host address, two services on the same host can bypass the lower layers of the TCP/IP stack. You can also ping the loopback address to test the configuration of TCP/IP on the local host.

Although only the single 127.0.0.1 address is used, addresses 127.0.0.0 to 127.255.255.255 are reserved. Any address within this block will loop back to the local host. No address within this block should ever appear on any network.

Link-Local Addresses

IPv4 addresses in the address block 169.254.0.0 to 169.254.255.255 (169.254.0.0/16) are designated as link-local addresses. These addresses can be automatically assigned to the local host by the operating system in environments where no IP configuration is available. These might be used in a small peer-to-peer network or for a host that could not automatically obtain an address from a DHCP server.

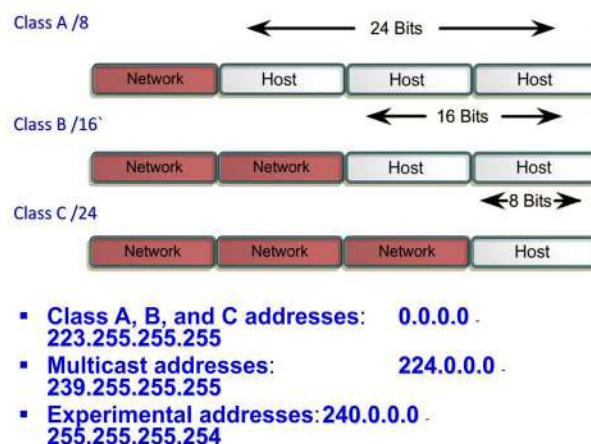
TEST-NET Addresses

The address block 192.0.2.0 to 192.0.2.255 (192.0.2.0/24) is set aside for teaching and learning purposes. These addresses can be used in documentation and network examples.

Experimental Addresses

The addresses in the block 240.0.0.0 to 255.255.255.254 are listed as reserved for future use (RFC 3330). Currently, these addresses can only be used for research or experimentation purposes, but cannot be used in an IPv4 network.

Legacy Classful Addresses



Assignment of IP Addresses



For a company or organization to have network hosts, such as web servers, accessible from the Internet, that organization must have a block of public addresses assigned. Remember that public addresses must be unique, and use of these public addresses is regulated and allocated to each organization separately. This is true for IPv4 and IPv6 addresses.

IANA and RIRs

Internet Assigned Numbers Authority (IANA) (<http://www.iana.org>) manages the allocation of IPv4 and IPv6 addresses. Until the mid-1990s, all IPv4 address space was managed directly by the IANA. At that time, the remaining IPv4 address space was allocated to various other registries to manage for particular purposes or for regional areas. These registration companies are called Regional Internet Registries (RIRs).

The major registries are:

AfrinIC (African Network Information Centre) - Africa Region <http://www.afrinic.net>

APNIC (Asia Pacific Network Information Centre) - Asia/Pacific Region <http://www.apnic.net>

ARIN (American Registry for Internet Numbers) - North America Region <http://www.arin.net>

LACNIC (Regional Latin-American and Caribbean IP Address Registry) - Latin America and some Caribbean Islands <http://www.lacnic.net>

RIPE NCC (Reseaux IP Europeans) - Europe, the Middle East, and Central Asia <http://www.ripe.net>

ISPs

RIRs are responsible for allocating IP addresses to the Internet Service Providers (ISPs). Most companies or organizations obtain their IPv4 address blocks from an ISP. An ISP will generally supply a small number of usable IPv4 addresses (6 or 14) to their customers as a part of their services. Larger blocks of addresses can be obtained based on justification of needs and for additional service costs.

Topic 29: Using Windows Calculator with Network Addresses:

In this activity, we will use windows calculator to calculate network addresses.

Topic 30: Converting IPv4 Addresses to Binary:

- Convert IPv4 Addresses from Dotted Decimal to Binary
- Bitwise ANDing
- Network Address Calculation

Topic 31: Network Segmentation:

In early network implementations, it was common for organizations to have all computers and other networked devices connected to a single IP network. All devices in the organization were assigned an IP address with a matching network ID. This type of configuration is known as a flat network design. In a small network, with a limited number of devices, a flat network design is not problematic. However, as the network grows, this type of configuration can create major issues.

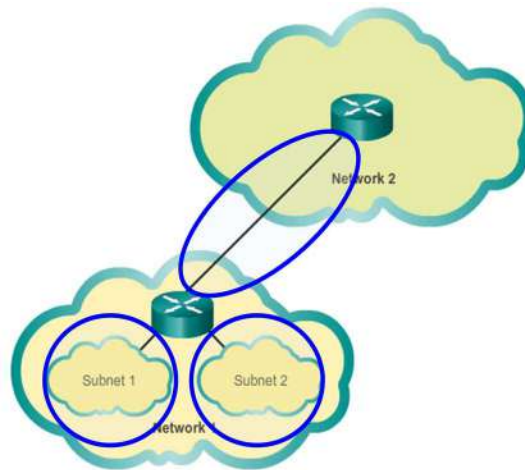
Consider how on an Ethernet LAN, devices use broadcasts to locate needed services and devices. Recall that a broadcast is sent to all hosts on an IP network. The Dynamic Host Configuration Protocol (DHCP) is an example of a network service that depends on broadcasts. Devices send broadcasts across the network to locate the DHCP server. On a large network, this could create a significant amount of traffic slowing network operations.

The process of segmenting a network, by dividing it into multiple smaller network spaces, is called subnetting. These sub-networks are called subnets. Network administrators can group devices and services into subnets that are determined by geographic location (perhaps the 3rd floor of a building), by organizational unit (perhaps the sales department), by device type (printers, servers, WAN), or any other division that makes sense for the network. Subnetting can reduce overall network traffic and improve network performance.

Subnetting

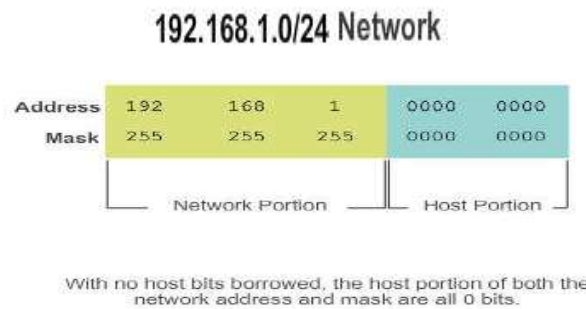
Segmenting networks in subnets creates smaller groups of devices and services in order to:

- Create smaller broadcast domains.
- Limit the amount of traffic on the other network segments.
- Provide low-level security.



- A router is required to subnet a network.
 - Each router interface is on a different subnet.
 - Devices on a subnet use the router interface as the default gateway.

Topic 32: Subnetting an IPv4 Network:



Every network address has a valid range of host addresses. All devices attached to the same network will have an IPv4 host address for that network and a common subnet mask or network prefix. The prefix and the subnet mask are different ways of representing the same thing - the network portion of an address.

IPv4 subnets are created by using one or more of the host bits as network bits. This is done by extending the mask to borrow some of the bits from the host portion of the address to create additional network bits. The more host bits borrowed, the more subnets that can be defined. For each bit borrowed, the number of subnetworks available is doubled. For example, if 1 bit is borrowed, 2 subnets can be created. If 2 bits, 4 subnets are created, if 3 bits are borrowed, 8 subnets are created, and so on. However, with each bit borrowed, fewer host addresses are available per subnet.

Bits can only be borrowed from the host portion of the address. The network portion of the address is allocated by the service provider and cannot be changed.

Topic 33: Calculating IPv4 Subnets:

- Calculate IPv4 Address Subnetting

Topic 34: Packet Tracer – Subnetting Scenario:

In this activity, you are given the network address of 192.168.100.0/24 to subnet and provide the IP addressing for the network shown in the topology. Each LAN in the network requires enough space for, at least, 25 addresses for end devices, the switch and the router. The connection between R1 to R2 will require an IP address for each end of the link.

- Design an IP Addressing Scheme
- Assign IP Addresses to Network Devices and Verify Connectivity

Topic 35: Packet Tracer – Subnetting Scenario - 2:

Part 2 of the previous activity.

Topic 36: Variable Length Subnet Masking (VLSM):

VLSM allows a network space to be divided in unequal parts. With VLSM the subnet mask will vary depending on how many bits have been borrowed for a particular subnet, thus the variable part of the VLSM.

VLSM subnetting is similar to traditional subnetting in that bits are borrowed to create subnets. The formulas to calculate the number of hosts per subnet and the number of subnets created still apply. The difference is that subnetting is not a single pass activity. With VLSM, the network is first subnetted, and then the subnets are subnetted again. This process can be repeated multiple times to create subnets of various sizes.

Please refer to the slides for more understanding of VLSM with the help of an example.

Topic 37: Anatomy of a Router:

How does clicking a link in a web browser return the desired information in mere seconds? Although there are many devices and technologies collaboratively working together to enable this, the primary device is the router. Stated simply, a router connects one network to another network.

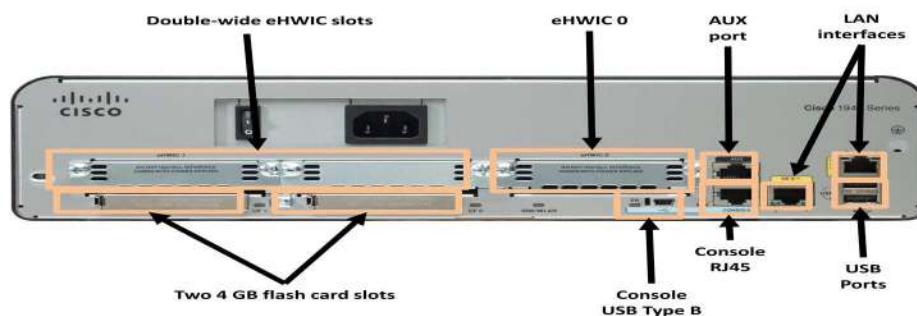
Communication between networks would not be possible without a router determining the best path to the destination and forwarding traffic to the next router along that path. The router is responsible for the routing of traffic between networks.

When a packet arrives on a router interface, the router uses its routing table to determine how to reach the destination network. The destination of the IP packet might be a web server in another country or an email server on the local area network. It is the responsibility of routers to deliver those packets efficiently. The effectiveness of internetwork communications depends, to a large degree, on the ability of routers to forward packets in the most efficient way possible.

Functions of a Router

- Routers are computers
- Routers interconnects networks
- Routers choose best paths
- A router has access to four types of memory: RAM, ROM, NVRAM, and Flash.
- RAM is used to store various applications and processes including:
- **Cisco IOS** - The IOS is copied into RAM during bootup.
- **Running configuration file** - This is the configuration file that stores the configuration commands that the router IOS is currently using. It is also known as the running-config.
- **IP routing table** - This file stores information about directly-connected and remote networks. It is used to determine the best path to use to forward packets.
- **ARP cache** - This cache contains the IPv4 address to MAC address mappings, similar to the Address Resolution Protocol (ARP) cache on a PC. The ARP cache is used on routers that have LAN interfaces, such as Ethernet interfaces.
- **Packet buffer** - Packets are temporarily stored in a buffer when received on an interface or before they exit an interface.

- Like computers, Cisco routers actually use dynamic random-access memory (DRAM). DRAM is a very common kind of RAM that stores the instructions and data needed to be executed by the CPU. Unlike ROM, RAM is volatile memory and requires continual power to maintain its information. It loses all of its content when the router is powered down or restarted.
- By default 1941 routers come with 512 MB of DRAM soldered on the main system board (onboard) and one dual in-line memory module (DIMM) slot for memory upgrades of up to an additional 2.0 GB. Cisco 2901, 2911, and 2921 models come with 512 MB of onboard DRAM. Note that first generation ISRs and older Cisco routers do not have onboard RAM.
- Cisco routers use ROM to store:
 - **Bootup instructions** - Provides the startup instructions.
 - **Basic diagnostic software** - Performs the power-on self-test (POST) of all components.
 - **Limited IOS** - Provides a limited backup version of the OS, in case the router cannot load the full featured IOS.
- ROM is firmware embedded on an integrated circuit inside the router and does not lose its contents when the router loses power or is restarted.
- NVRAM is used by the Cisco IOS as permanent storage for the startup configuration file (startup-config). Like ROM, NVRAM does not lose its contents when power is turned off.
- Flash memory is non-volatile computer memory used as permanent storage for the IOS and other system related files. The IOS is copied from flash into RAM during the bootup process.
- Cisco 1941 routers come with two external Compact Flash slots. Each slot can support high-speed storage densities upgradeable to 4GB in density.



A Cisco 1941 router includes the following connections:

Console ports - Two console ports for the initial configuration and command-line interface (CLI) management access using a regular RJ-45 port and a new USB Type-B (mini-B USB) connector.

AUX port - An RJ-45 port for remote management access; this is similar to the Console port.

Two LAN interfaces - Two Gigabit Ethernet interfaces for LAN access.

Enhanced high-speed WAN interface card (EHWIC) slots - Two slots that provide modularity and flexibility by enabling the router to support different types of interface modules, including Serial, digital subscriber line (DSL), switch port, and wireless.

The Cisco 1941 ISR also has storage slots to support expanded capabilities. Dual-compact flash memory slots are capable of supporting a 4 GB compact flash card each for increased storage space. Two USB host ports are included for additional storage space and secure token capability.

Compact flash can store the Cisco IOS software image, log files, voice configuration files, HTML files, backup configurations, or any other file needed for the system. By default, only slot 0 is populated with a compact flash card from the factory, and it is the default boot location.

The figure above identifies the location of these connections and slots.

Cisco devices, routers, and switches typically interconnect many devices. For this reason, these devices have several types of ports and interfaces. These ports and interfaces are used to connect cables to the device.

The connections on a Cisco router can be grouped into two categories:

Management ports - These are the console and auxiliary ports used to configure, manage, and troubleshoot the router. Unlike LAN and WAN interfaces, management ports are not used for packet forwarding.

Inband Router interfaces - These are the LAN and WAN interfaces configured with IP addressing to carry user traffic. Ethernet interfaces are the most common LAN connections, while common WAN connections include serial and DSL interfaces.

Router Interfaces

Similar to a Cisco switch, there are several ways to access the CLI environment on a Cisco router. The most common methods are:

Console - Uses a low speed serial or USB connection to provide direct connect, out-of-band management access to a Cisco device.

Telnet or SSH - Two methods for remotely accessing a CLI session across an active network interface.

AUX port - Used for remote management of the router using a dial-up telephone line and modem.

The console and AUX port are located on the router.

In addition to these ports, routers also have network interfaces to receive and forward IP packets. Routers have multiple interfaces that are used to connect to multiple networks. Typically, the interfaces connect to various types of networks, which mean that different types of media and connectors are required.

Every interface on the router is a member or host on a different IP network. Each interface must be configured with an IP address and subnet mask of a different network. The Cisco IOS does not allow two active interfaces on the same router to belong to the same network.

Router interfaces can be grouped into two categories:

Ethernet LAN interfaces - Used for connecting cables that terminate with LAN devices, such as computers and switches. This interface can also be used to connect routers to each other. Several conventions for naming Ethernet interfaces are popular: the older Ethernet, FastEthernet, and Gigabit Ethernet. The name used depends on the device type and model.

Serial WAN interfaces - Used for connecting routers to external networks, usually over a larger geographical distance. Similar to LAN interfaces, each serial WAN interface has its own IP address and subnet mask, which identifies it as a member of a specific network.

Topic 38: Packet Tracer – Exploring Internetworking Devices:

- Identify Physical Characteristics of Internetworking Devices
- Select Correct Modules for Connectivity
- Connect Devices

In this activity, you will explore the different options available on internetworking devices. You will also be required to determine which options provide the necessary connectivity when connecting multiple devices. Finally, you will add the correct modules and connect the devices.

Topic 39: Router Bootup:

The Cisco IOS operational details vary on different internetworking devices, depending on the device's purpose and feature set. However, Cisco IOS for routers provides the following:

- Addressing
- Interfaces
- Routing
- Security
- QoS
- Resources Management

The IOS file itself is several megabytes in size and similar to Cisco IOS switches, is stored in flash memory. Using flash allows the IOS to be upgraded to newer versions or to have new features added. During bootup, the IOS is copied from flash memory into RAM. DRAM is much faster than flash; therefore, copying the IOS into RAM increases the performance of the device.

A router loads the following two files into RAM when it is booted:

IOS image file - The IOS facilitates the basic operation of the device's hardware components. The IOS image file is stored in flash memory.

Startup configuration file - The startup configuration file contains commands that are used to initially configure a router and create the running configuration file stored in RAM. The startup configuration file is stored in NVRAM. All configuration changes are stored in the running configuration file and are implemented immediately by the IOS.

The running configuration is modified when the network administrator performs device configuration. When changes are made to the running-config file, it should be saved to NVRAM as the startup configuration file, in case the router is restarted or loses power.

Topic 40: Configuring Routers:

Cisco routers and Cisco switches have many similarities. They support a similar modal operating system, similar command structures, and many of the same commands. In addition, both devices have similar initial configuration steps.

When configuring a Cisco switch or router, the following basic tasks should be performed first:

Name the device - Distinguishes it from other routers.

Secure management access - Secures privileged EXEC, user EXEC, and Telnet access, and encrypts passwords to their highest level.

Configure a banner - Provides legal notification of unauthorized access.

One distinguishing feature between switches and routers is the type of interfaces supported by each. For example, Layer 2 switches support LANs and, therefore, have multiple FastEthernet or Gigabit Ethernet ports.

Routers support LANs and WANs and can interconnect different types of networks; therefore, they support many types of interfaces. For example, G2 ISRs have one or two integrated Gigabit Ethernet interfaces and High-Speed WAN Interface Card (HWIC) slots to accommodate other types of network interfaces, including serial, DSL, and cable interfaces.

To be available, an interface must be:

If using IPv4, configured with an address and a subnet mask - Use the `ip address ip-address subnet-mask` interface configuration command.

Activated - By default, LAN and WAN interfaces are not activated (**shutdown**). To enable an interface, it must be activated using the **no shutdown** command. (This is similar to powering on the interface.) The interface must also be connected to another device (a hub, a switch, or another router) for the physical layer to be active.

Optionally, the interface could also be configured with a short description. It is good practice to configure a description on each interface. The description text is limited to 240 characters. On production networks, a description can be helpful in troubleshooting by providing information about the type of network to which the interface is connected. If the interface connects to an ISP or service carrier, it is helpful to enter the third party connection and contact information.

Topic 41: Packet Tracer – Configure Initial Router Settings:

- Verify the Default Router Configuration
- Verify and Configure Initial Router Configuration
- Save the Running Configuration File

In this activity, you will perform basic router configurations. You will secure access to the CLI and console port using encrypted and plain text passwords. You will also configure messages for users logging into the router. These banners also warn unauthorized users that access is prohibited. Finally, you will verify and save your running configuration.

Topic 42: Packet Tracer – Configuring IPv4:

- Configure IPv4 Addressing and Verify Connectivity

Topic 43: Verify Connectivity of Directly Connected Networks:

There are several **show** commands that can be used to verify the operation and configuration of an interface. The following three commands are especially useful to quickly identify an interface status:

show ip interface brief - Displays a summary for all interfaces including the IPv4 address of the interface and current operational status.

show ip route - Displays the contents of the IPv4 routing table stored in RAM. In Cisco IOS 15, active interfaces should appear in the routing table with two related entries identified by the code C (Connected) or L (Local).

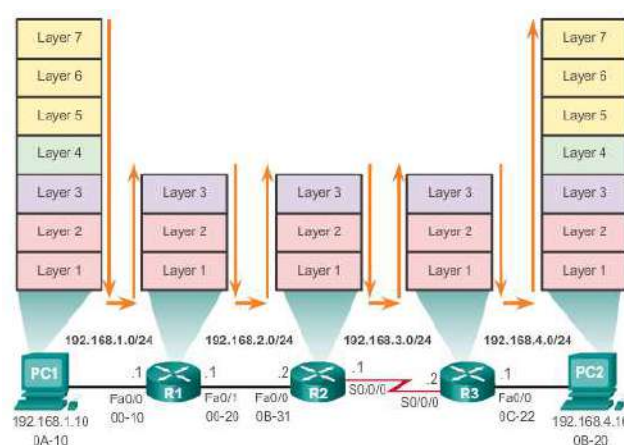
show running-config interface interface-id - Displays the commands configured on the specified interface.

The following two commands are used to gather more detailed interface information:

show interfaces - Displays interface information and packet flow count for all interfaces on the device.

show ip interface - Displays the IPv4 related information for all interfaces on a router.

Topic 44: Switching Packets Between Networks:



A primary function of a router is to forward packets toward their destination. This is accomplished by using a switching function, which is the process used by a router to accept a packet on one interface and forward it out of another interface. A key responsibility of the switching function is to encapsulate packets in the appropriate data link frame type for the outgoing data link.

After the router has determined the exit interface using the path determination function, the router must encapsulate the packet into the data link frame of the outgoing interface.

What does a router do with a packet received from one network and destined for another network? The router performs the following three major steps:

Step 1. De-encapsulates the Layer 3 packet by removing the Layer 2 frame header and trailer.

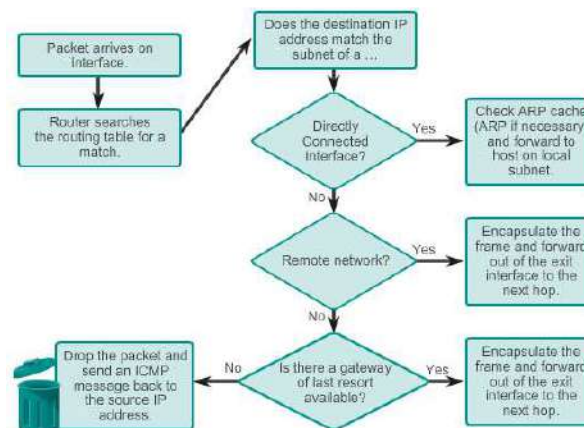
Step 2. Examines the destination IP address of the IP packet to find the best path in the routing table.

Step 3. If the router finds a path to the destination, it encapsulates the Layer 3 packet into a new Layer 2 frame and forwards the frame out the exit interface.

As shown in the figure above, devices have Layer 3 IPv4 addresses and Ethernet interfaces have Layer 2 data link addresses. For example, PC1 is configured with IPv4 address 192.168.1.10 and an example MAC address of 0A-10. As a packet travels from the source device to the final destination device, the Layer 3 IP addresses do not change. However, the Layer 2 data link addresses change at every hop as the packet is de-encapsulated and re-encapsulated in a new frame by each router. It is very likely that the packet is encapsulated in a different type of Layer 2 frame than the one in which it was received. For example, an Ethernet encapsulated frame might be received by the router on a FastEthernet interface, and then processed to be forwarded out of a serial interface as a Point-to-Point Protocol (PPP) encapsulated frame.

Please refer to the slides for detail understanding with the help of an example.

Topic 45: Path Determination:



A primary function of a router is to determine the best path to use to send packets. To determine the best path, the router searches its routing table for a network address that matches the destination IP address of the packet.

The routing table search results in one of three path determinations:

Directly connected network - If the destination IP address of the packet belongs to a device on a network that is directly connected to one of the interfaces of the router, that packet is forwarded directly to the destination device. This means that the destination IP address of the packet is a host address on the same network as the interface of the router.

Remote network - If the destination IP address of the packet belongs to a remote network, then the packet is forwarded to another router. Remote networks can only be reached by forwarding packets to another router.

No route determined - If the destination IP address of the packet does not belong to either a connected or remote network, the router determines if there is a Gateway of Last Resort available. A Gateway of Last Resort is set when a default route is configured on a router. If there is a default route, the packet is forwarded to the Gateway of Last Resort. If the router does not have a default route, then the packet is discarded. If the packet is discarded, the router sends an ICMP unreachable message to the source IP address of the packet.

The logic flowchart in the figure above illustrates the router packet forwarding decision process.

Determining the best path involves the evaluation of multiple paths to the same destination network and selecting the optimum or shortest path to reach that network. Whenever multiple paths to the same network exist, each path uses a different exit interface on the router to reach that network.

The best path is selected by a routing protocol based on the value or metric it uses to determine the distance to reach a network. A metric is the quantitative value used to measure the distance to a given network. The best path to a network is the path with the lowest metric.

Dynamic routing protocols typically use their own rules and metrics to build and update routing tables. The routing algorithm generates a value, or a metric, for each path through the network. Metrics can be based on either a single characteristic or several characteristics of a path. Some routing protocols can base route selection on multiple metrics, combining them into a single metric.

The following lists some dynamic protocols and the metrics they use:

Routing Information Protocol (RIP) - Hop count

Open Shortest Path First (OSPF) - Cisco's cost based on cumulative bandwidth from source to destination

Enhanced Interior Gateway Routing Protocol (EIGRP) - Bandwidth, delay, load, reliability

Load Balancing

What happens if a routing table has two or more paths with identical metrics to the same destination network?

When a router has two or more paths to a destination with equal cost metrics, then the router forwards the packets using both paths equally. This is called equal cost load balancing. The routing table contains the single destination network, but has multiple exit interfaces, one for each equal cost path. The router forwards packets using the multiple exit interfaces listed in the routing table.

If configured correctly, load balancing can increase the effectiveness and performance of the network. Equal cost load balancing can be configured to use both dynamic routing protocols and static routes.

Note: Only EIGRP supports unequal cost load balancing.

Topic 46: Analyze the Routing Table:

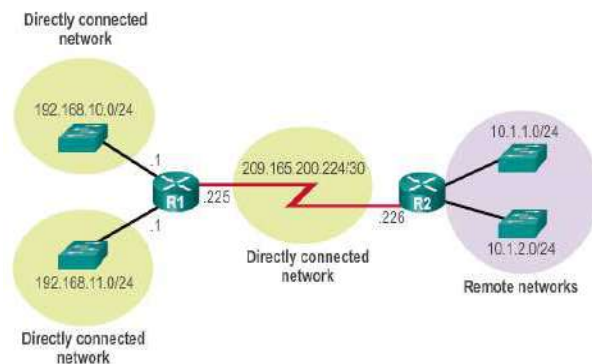
The routing table of a router stores information about:

Directly connected routes - These routes come from the active router interfaces. Routers add a directly connected route when an interface is configured with an IP address and is activated.

Remote routes - These are remote networks connected to other routers. Routes to these networks can either be statically configured or dynamically configured using dynamic routing protocols.

Specifically, a routing table is a data file in RAM that is used to store route information about directly connected and remote networks. The routing table contains network or next hop associations. These associations tell a router that a particular destination can be optimally reached by sending the packet to a specific router that represents the next hop on the way to the final destination. The next hop association can also be the outgoing or exit interface to the next destination.

The figure below identifies the directly connected networks and remote networks of router R1.



On a Cisco IOS router, the **show ip route** command can be used to display the IPv4 routing table of a router. A router provides additional route information, including how the route was learned, how long the route has been in the table, and which specific interface to use to get to a predefined destination.

Entries in the routing table can be added as:

Local Route interfaces - Added when an interface is configured and active. This entry is only displayed in IOS 15 or newer for IPv4 routes and all IOS releases for IPv6 routes.

Directly connected interfaces - Added to the routing table when an interface is configured and active.

Static routes - Added when a route is manually configured and the exit interface is active.

Dynamic routing protocol - Added when routing protocols that dynamically learn about the network, such as EIGRP or OSPF, are implemented and networks are identified.

The sources of the routing table entries are identified by a code. The code identifies how the route was learned. For instance, common codes include:

L - Identifies the address assigned to a router's interface. This allows the router to efficiently determine when it receives a packet for the interface instead of being forwarded.

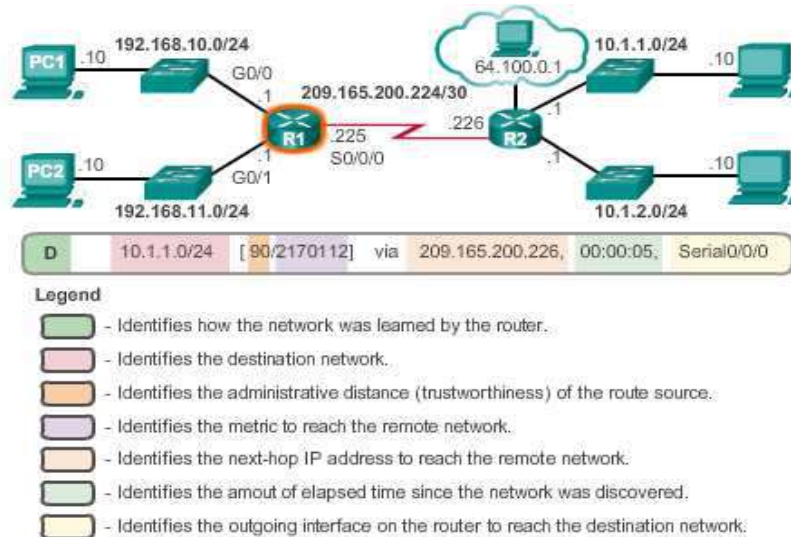
C - Identifies a directly connected network.

S - Identifies a static route created to reach a specific network.

D - Identifies a dynamically learned network from another router using EIGRP.

O - Identifies a dynamically learned network from another router using the OSPF routing protocol.

The figure below shows the routing table of R1 in a simple network.



Topic 47: Directly Connected/Static/Dynamic Routes:

Directly Connected Routes

An active, properly configured, directly connected interface actually creates two routing table entries. The figure below displays the IPv4 routing table entries on R1 for the directly connected network 192.168.10.0.

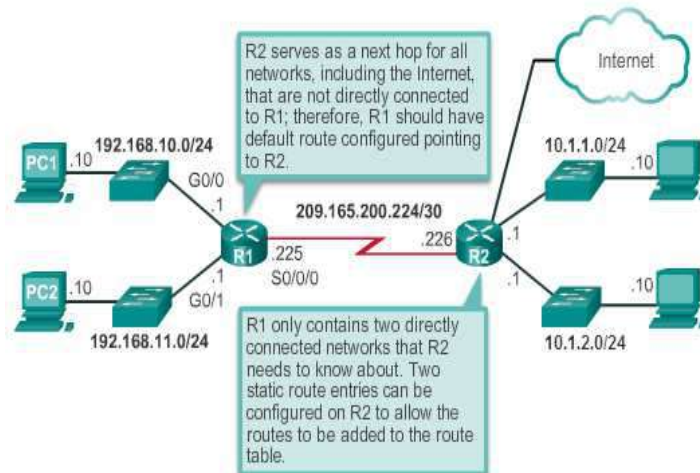
The routing table entry for directly connected interfaces is simpler than the entries for remote networks. The entries contain the following information:

Route source - Identifies how the route was learned. Directly connected interfaces have two route source codes. C identifies a directly connected network. L identifies the IPv4 address assigned to the router's interface.

Destination network - The address of the remote network.

Outgoing interface - Identifies the exit interface to use when forwarding packets to the destination network.

Note: Prior to IOS 15, local route routing table entries (L) were not displayed in the IPv4 routing table. Local route (L) entries have always been a part of the IPv6 routing table.



After directly connected interfaces are configured and added to the routing table, then static or dynamic routing can be implemented.

Static routes are manually configured. They define an explicit path between two networking devices. Unlike a dynamic routing protocol, static routes are not automatically updated and must be manually reconfigured if the network topology changes. The benefits of using static routes include improved security and resource efficiency. Static routes use less bandwidth than dynamic routing protocols, and no CPU cycles are used to calculate and communicate routes. The main disadvantage to using static routes is the lack of automatic reconfiguration if the network topology changes.

There are two common types of static routes in the routing table:

- Static route to a specific network
- Default static route

A static route can be configured to reach a specific remote network. IPv4 static routes are configured using the **ip route network mask {next-hop-ip | exit-intf}** global configuration command. A static route is identified in the routing table with the code S.

A default static route is similar to a default gateway on a host. The default static route specifies the exit point to use when the routing table does not contain a path for the destination network.

A default static route is useful when a router has only one exit point to another router, such as when the router connects to a central router or service provider.

To configure an IPv4 default static route, use the **ip route 0.0.0.0 0.0.0.0{exit-intf | next-hop-ip}** global configuration command.

Dynamic Routing

Dynamic routing protocols are used by routers to share information about the reachability and status of remote networks. Dynamic routing protocols perform several activities, including network discovery and maintaining routing tables.

Network discovery is the ability of a routing protocol to share information about the networks that it knows about with other routers that are also using the same routing protocol. Instead of depending on manually configured static routes to remote networks on every router, a dynamic routing protocol allows the routers to automatically learn about these networks from other routers. These networks, and the best path to each, are added to the routing table of the router, and identified as a network learned by a specific dynamic routing protocol.

During network discovery, routers exchange routes and update their routing tables. Routers have converged after they have finished exchanging and updating their routing tables. Routers then maintain the networks in their routing tables.

A router running a dynamic routing protocol does not only make a best path determination to a network, it also determines a new best path if the initial path becomes unusable (or if the topology changes). For these reasons, dynamic routing protocols have an advantage over static routes. Routers that use dynamic routing protocols automatically share routing information with other routers and compensate for any topology changes without involving the network administrator.

Cisco ISR routers can support a variety of dynamic IPv4 routing protocols including:

EIGRP - Enhanced Interior Gateway Routing Protocol

OSPF - Open Shortest Path First

IS-IS - Intermediate System-to-Intermediate System

RIP - Routing Information Protocol

To determine which routing protocols are supported by the IOS, use the **router ?** command in global configuration mode.

Topic 48: Packet Tracer – Configuring and Verifying a Small Network:

- Configure Devices and Verify Connectivity Gather Information with Show Commands

In this activity, you will configure a router with basic settings including IP addressing. You will also configure a switch for remote management and configure the PCs. After you have successfully verified connectivity, you will use **show** commands to gather information about the network.

Topic 49: Packet Tracer – Configuring & Verifying a Small Network - 2:

This activity is part 2 of the previous activity.

Topic 50: Packet Tracer – Connect a Router to a LAN:

- Display Router Information
- Configure Router Interfaces
- Verify the Configuration

In this activity, you will use various **show** commands to display the current state of the router. You will then use the Addressing Table to configure router Ethernet interfaces. Finally, you will use commands to verify and test your configurations.

Topic 51: Testing the Network: Ping and ICMPv4:

Although IP is not a reliable protocol, the TCP/IP suite does provide for messages to be sent in the event of certain errors. These messages are sent using the services of ICMP. The purpose of these messages is to provide feedback about issues related to the processing of IP packets under certain conditions, not to make IP reliable. ICMP messages are not required and are often not allowed within a network for security reasons.

ICMP is available for both IPv4 and IPv6. ICMPv4 is the messaging protocol for IPv4. ICMPv6 provides these same services for IPv6 but includes additional functionality. In this course, the term ICMP will be used when referring to both ICMPv4 and ICMPv6.

ICMP messages common to both ICMPv4 and ICMPv6 include:

- Host confirmation
- Destination or Service Unreachable
- Time exceeded
- Route redirection

Host Confirmation

An ICMP Echo Message can be used to determine if a host is operational. The local host sends an ICMP Echo Request to a host. If the host is available, the destination host responds with an Echo Reply.

Destination or Service Unreachable

When a host or gateway receives a packet that it cannot deliver, it can use an ICMP Destination Unreachable message to notify the source that the destination or service is unreachable. The message will include a code that indicates why the packet could not be delivered.

Some of the Destination Unreachable codes for ICMPv4 are:

- 0 - net unreachable.
- 1 - host unreachable.
- 2 - protocol unreachable.
- 3 - port unreachable.

Time Exceeded

An ICMPv4 Time Exceeded message is used by a router to indicate that a packet cannot be forwarded because the Time to Live (TTL) field of the packet was decremented to 0. If a router receives a packet and decrements the TTL field in the IPv4 packet to zero, it discards the packet and sends a Time Exceeded message to the source host.

ICMPv6 also sends a Time Exceeded message if the router cannot forward an IPv6 packet because the packet has expired. IPv6 does not have a TTL field; it uses the hop limit field to determine if the packet has expired.

Route Redirection

A router may use the ICMP Redirect Message to notify the hosts on a network that a better route is available for a particular destination. This message may only be used when the source host is on the same physical network as both gateways.

Both ICMPv4 and ICMPv6 use route redirection messages.

Ping

Ping is a testing utility that uses ICMP echo request and echo reply messages to test connectivity between hosts. Ping works with both IPv4 and IPv6 hosts.

To test connectivity to another host on a network, an echo request is sent to the host address using the **ping** command. If the host at the specified address receives the echo request, it responds with an echo reply. As each echo reply is received, ping provides feedback on the time between when the request was sent and when the reply was received. This can be a measure of network performance.

Ping has a timeout value for the reply. If a reply is not received within the timeout, ping provides a message indicating that a response was not received. This usually indicates that there is a problem, but could also indicate that security features blocking ping messages have been enabled on the network.

After all the requests are sent, the ping utility provides a summary that includes the success rate and average round-trip time to the destination.

You can also use ping to test the ability of a host to communicate on the local network. This is generally done by pinging the IP address of the gateway of the host. A ping to the gateway indicates that the host and the router interface serving as the gateway are both operational on the local network.

Ping can also be used to test the ability of a local host to communicate across an internetwork. The local host can ping an operational IPv4 host of a remote network.

If this ping is successful, the operation of a large piece of the internetwork can be verified. A successful ping across the internetwork confirms communication on the local network, the operation of the router serving as our gateway, and the operation of all other routers that might be in the path between the local network and the network of the remote host.

Traceroute

Ping is used to test connectivity between two hosts, but does not provide information about the details of devices between the hosts. Traceroute (tracert) is a utility that generates a list of hops that were successfully reached along the path. This list can provide important verification and troubleshooting information. If the data reaches the destination, then the trace lists the interface of every router in the path between the hosts. If the data fails at some hop along the way, the address of the last router that responded to the trace can provide an indication of where the problem or security restrictions are found.

Round Trip Time (RTT)

Using traceroute provides round trip time for each hop along the path and indicates if a hop fails to respond. The round trip time is the time a packet takes to reach the remote host and for the response from the host to return. An asterisk (*) is used to indicate a lost or unreplied packet.

This information can be used to locate a problematic router in the path. If the display shows high response times or data losses from a particular hop, this is an indication that the resources of the router or its connections may be stressed.

IPv4 Time-to-Live (TTL) and IPv6 Hop Limit

Traceroute makes use of a function of the TTL field in IPv4 and the Hop Limit field in IPv6 in the Layer 3 headers, along with the ICMP time exceeded message.

The first sequence of messages sent from traceroute will have a TTL field value of 1. This causes the TTL to time out the IPv4 packet at the first router. This router then responds with an ICMPv4 message. Traceroute now has the address of the first hop.

Traceroute then progressively increments the TTL field (2, 3, 4...) for each sequence of messages. This provides the trace with the address of each hop as the packets timeout further down the path. The TTL field continues to be increased until the destination is reached or it is incremented to a predefined maximum.

Once the final destination is reached, the host responds with either an ICMP port unreachable message or an ICMP echo reply message instead of the ICMP time exceeded message.

Topic 52: Packet Tracer – Building a Switch and Router Network - 1:

- Setup Topology
- Configure Devices
- Verify Connectivity
- Display Device Information

In this activity, you will configure a router with basic settings including IP addressing. You will also configure a switch for remote management and configure the PCs. After you have successfully verified connectivity, you will use **show** commands to gather information about the network.

Topic 53: Packet Tracer – Building a Switch and Router Network - 2:

This is the continuation of the previous topic.

Topic 54: Packet Tracer – Testing Network Connectivity with Ping & Traceroute:

- Build and Configure a Network
- Ping Command
- Tracert/Traceroute Command

In this activity, you will configure a router with basic settings including IP addressing. You will also configure a switch for remote management and configure the PCs. After you have successfully verified connectivity, you will use show commands to gather information about the network.

Topic 55: Packet Tracer – Testing Network Connectivity with Ping & Traceroute - 2:

This is the continuation of the previous topic.

Topic 56: Packet Tracer – Testing Network Connectivity with Ping & Traceroute - 3:

This is the continuation of the previous topic.

Topic 57: Static Routing:

A router can learn about remote networks in one of two ways:

Manually - Remote networks are manually entered into the route table using static routes.

Dynamically - Remote routes are automatically learned using a dynamic routing protocol.

A network administrator can manually configure a static route to reach a specific network. Unlike a dynamic routing protocol, static routes are not automatically updated and must be manually reconfigured any time the network topology changes. A static route does not change until the administrator manually reconfigures it.

Dynamic vs. Static Routing

Static routing has three primary uses:

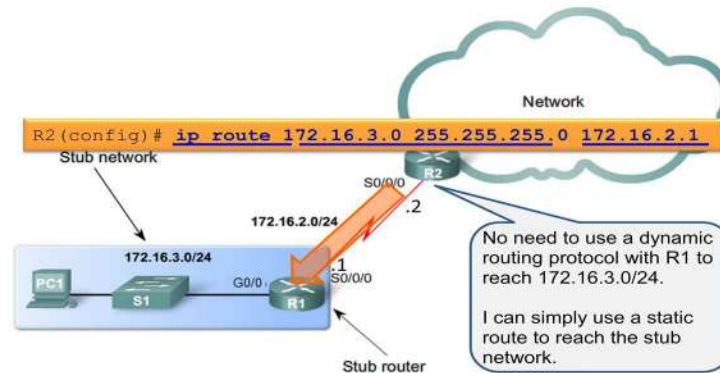
- Providing ease of routing table maintenance in smaller networks that are not expected to grow significantly.
- Routing to and from stub networks. A stub network is a network accessed by a single route, and the router has only one neighbor.
- Using a single default route to represent a path to any network that does not have a more specific match with another route in the routing table. Default routes are used to send traffic to any destination beyond the next upstream router.

Static routes can also be used to:

- Reduce the number of routes advertised by summarizing several contiguous networks as one static route
- Create a backup route in case a primary route link fails

The following types of IPv4 and IPv6 static routes will be discussed:

- Standard static route
- Default static route
- Summary static route
- Floating static route



A default static route is a route that matches all packets. A default route identifies the gateway IP address to which the router sends all IP packets that it does not have a learned or static route. A default static route is simply a static route with 0.0.0.0/0 as the destination IPv4 address. Configuring a default static route creates a Gateway of Last Resort.

Default static routes are used:

- When no other routes in the routing table match the packet destination IP address. In other words, when a more specific match does not exist. A common use is when connecting a company's edge router to the ISP network.
- When a router has only one other router to which it is connected. This condition is known as a stub router.

Summary Static Route

To reduce the number of routing table entries, multiple static routes can be summarized into a single static route if:

- The destination networks are contiguous and can be summarized into a single network address.
- The multiple static routes all use the same exit interface or next-hop IP address.

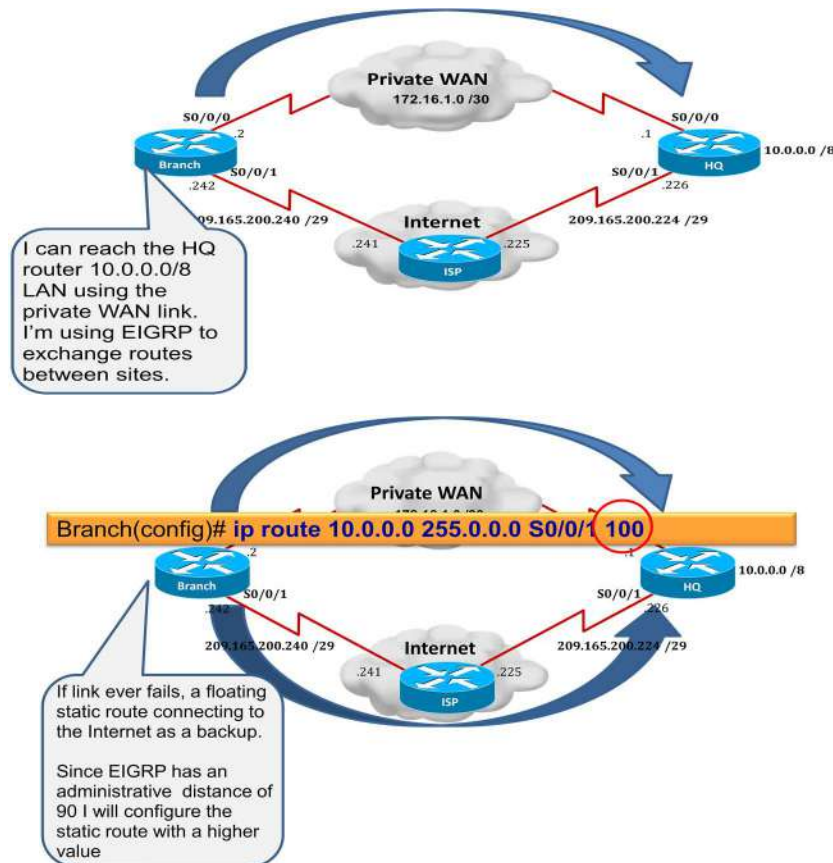
Floating Static Route

Another type of static route is a floating static route. Floating static routes are static routes that are used to provide a backup path to a primary static or dynamic route, in the event of a link failure. The floating static route is only used when the primary route is not available.

To accomplish this, the floating static route is configured with a higher administrative distance than the primary route. Recall that the administrative distance represents the trustworthiness of a route. If multiple paths to the destination exist, the router will choose the path with the lowest administrative distance.

For example, assume that an administrator wants to create a floating static route as a backup to an EIGRP-learned route. The floating static route must be configured with a higher administrative distance than EIGRP. EIGRP has an administrative distance of 90. If the floating static route is configured with an administrative distance of 95, the dynamic route learned through EIGRP is preferred to the floating static route. If the EIGRP-learned route is lost, the floating static route is used in its place.

In the figure below, the Branch router typically forwards all traffic to the HQ router over the private WAN link. In this example, the routers exchange route information using EIGRP. A floating static route, with an administrative distance of 91 or higher, could be configured to serve as a backup route. If the private WAN link fails and the EIGRP route disappears from the routing table, the router selects the floating static route as the best path to reach the HQ LAN.



Topic 58: Static Routing - 2:

This is the continuation of the previous topic.

Topic 59: Configure Static & Default Routes:

ip route network-add subnet {ip-address | exit-intf [ip-address]} [distance]

Parameter	Description
<i>network-add</i>	• Destination network address of the remote network to be added to the routing table.
<i>subnet</i>	• Subnet mask of the remote network to be added to the routing table. • Note: The subnet mask can be modified to summarize a group of networks
<i>ip-address</i>	• Commonly referred to as the next-hop router's IP address.
<i>exit-intf</i>	• Use the outgoing interface to forward packets to the destination network. • Also referred to as a directly attached static route. • Typically used when connecting in a point-to-point configuration.
<i>distance</i>	• Used to create a floating static route by setting an administrative distance that is higher than a dynamically learned route.

Next Hop Static Route

In a next-hop static route, only the next-hop IP address is specified. The output interface is derived from the next hop.

Please see the slide for an example.

Directly Connected Static Route

When configuring a static route, another option is to use the exit interface to specify the next-hop address. In older IOS versions, prior to CEF, this method is used to avoid the recursive lookup problem.

Configuring a directly connected static route with an exit interface allows the routing table to resolve the exit interface in a single search, instead of two searches.

Note: For point-to-point interfaces, you can use static routes that point to the exit interface or to the next-hop address. For multipoint/broadcast interfaces, it is more suitable to use static routes that point to a next-hop address.

Please see the slide for an example.

Fully Specified Static Route

In a fully specified static route, both the output interface and the next-hop IP address are specified. This is another type of static route that is used in older IOSs, prior to CEF. This form of static route is used when the output interface is a multi-access interface and it is necessary to explicitly identify the next hop. The next hop must be directly connected to the specified exit interface.

Please see the slide for an example.

Verify a Static Route

R1# show ip route static | begin Gateway

Gateway of last resort is not set

172.16.0.0/16 is variably subnetted, 5 subnets, 2 masks

S 172.16.1.0/24 [1/0] via 172.16.2.2

S 192.168.1.0/24 [1/0] via 172.16.2.2

S 192.168.2.0/24 [1/0] via 172.16.2.2

R1# show ip route 192.168.2.1

Routing entry for 192.168.2.0/24

Known via "static", distance 1, metric 0

Routing Descriptor Blocks:

* 172.16.2.2

Route metric is 0, traffic share count is 1

R1# show running-config | section ip route

ip route 172.16.1.0 255.255.255.0 172.16.2.2

ip route 192.168.1.0 255.255.255.0 172.16.2.2

ip route 192.168.2.0 255.255.255.0 172.16.2.2

Default Static Route

```
ip route 0.0.0.0 0.0.0.0 {ip-address | exit-intf [ip-address]}
```

Parameter	Description
0.0.0.0	• Matches any network address.
0.0.0.0	• Matches any subnet mask.
<i>ip-address</i>	• Commonly referred to as the next-hop router's IP address. • Typically used when connecting to a broadcast media (i.e., Ethernet) . • Commonly creates a recursive lookup.
<i>exit-intf</i>	• Use the outgoing interface to forward packets to the destination network. • Also referred to as a directly attached static route. • Typically used when connecting in a point-to-point configuration.

Configure Default Static Route

A default route is a static route that matches all packets. Rather than storing all routes to all networks in the routing table, a router can store a single default route to represent any network that is not in the routing table.

Routers commonly use default routes that are either configured locally or learned from another router, using a dynamic routing protocol. A default route is used when no other routes in the routing table match the destination IP address of the packet. In other words, if a more specific match does not exist, then the default route is used as the Gateway of Last Resort.

Default static routes are commonly used when connecting:

- An edge router to a service provider network

- A stub router (a router with only one upstream neighbor router)

The command syntax for a default static route is similar to any other static route, except that the network address is **0.0.0.0** and the subnet mask is **0.0.0.0**. The basic command syntax of a default static route is:

ip route 0.0.0.0 0.0.0.0 {ip-address | exit-intf }

Verify Default Static Route

```
R1# show ip route static
Codes: L - local, C - connected, S - static, R - RIP, M - mobile, B - BGP
       D - EIGRP, EX - EIGRP external, O - OSPF, IA - OSPF inter area
       N1 - OSPF NSSA external type 1, N2 - OSPF NSSA external type 2
       E1 - OSPF external type 1, E2 - OSPF external type 2
       i - IS-IS, su - IS-IS summary, L1 - IS-IS level-1, L2 - IS-IS level-2
       ia - IS-IS inter area, * - candidate default, U - per-user static route
       o - ODR, P - periodic downloaded static route, H - NHRP, I - LISP
       + - replicated route, % - next hop override

Gateway of last resort is 0.0.0.0 to network 0.0.0.0

S*   0.0.0.0/0 is via 172.16.2.2
R1#
```

Topic 60: Configure Static & Default Routes - 2:

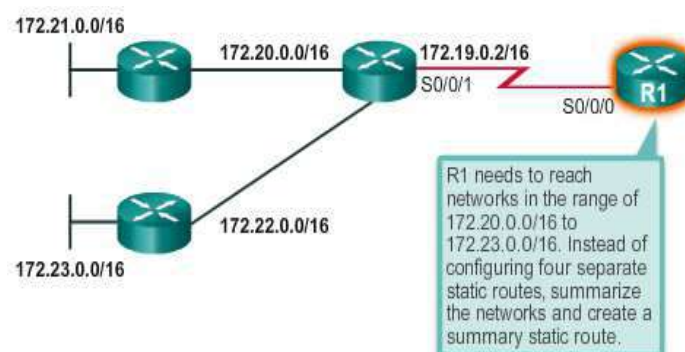
This is the continuation of the previous topic.

Topic 61: Configure Summary Static Routes:

Route summarization, also known as route aggregation, is the process of advertising a contiguous set of addresses as a single address with a less-specific, shorter subnet mask. CIDR is a form of route summarization and is synonymous with the term supernetting.

CIDR ignores the limitation of classful boundaries, and allows summarization with masks that are smaller than that of the default classful mask. This type of summarization helps reduce the number of entries in routing updates and lowers the number of entries in local routing tables. It also helps reduce bandwidth utilization for routing updates and results in faster routing table lookups.

In the figure below, R1 requires a summary static route to reach networks in the range of 172.20.0.0/16 to 172.23.0.0/16.



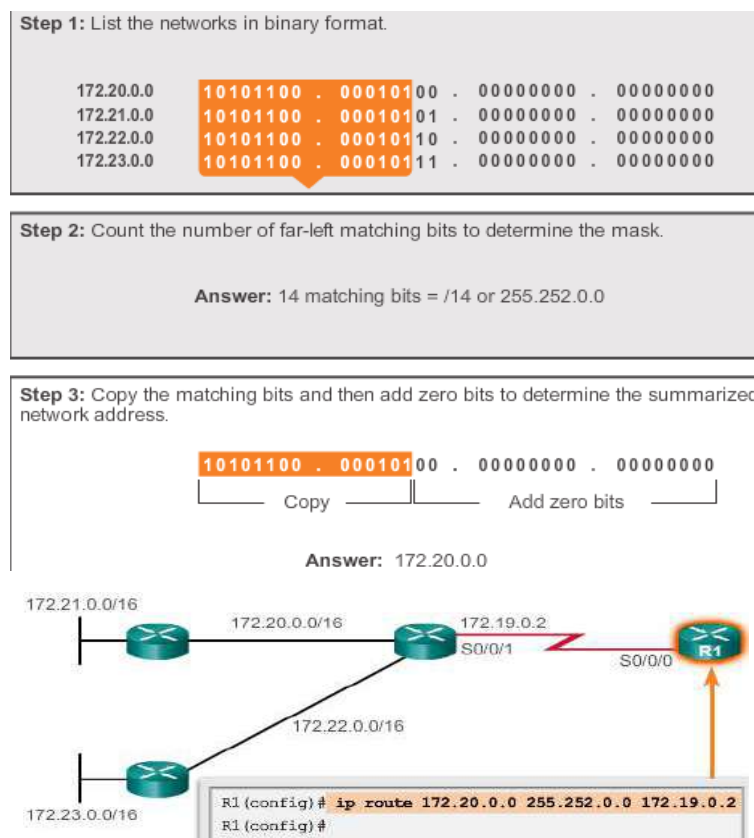
Summarizing networks into a single address and mask can be done in three steps:

Step 1. List the networks in binary format. Figure below lists networks 172.20.0.0/16 to 172.23.0.0/16 in binary format.

Step 2. Count the number of far-left matching bits to determine the mask for the summary route. Figure below highlights the 14 far left matching bits. This is the prefix, or subnet mask, for the summarized route: /14 or 255.252.0.0.

Step 3. Copy the matching bits and then add zero bits to determine the summarized network address. Figure below shows that the matching bits with zeros at the end results in the network address 172.20.0.0. The four networks - 172.20.0.0/16, 172.21.0.0/16, 172.22.0.0/16, and 172.23.0.0/16 - can be summarized into the single network address and prefix 172.20.0.0/14.

Figure below displays R1 configured with a summary static route to reach networks 172.20.0.0/16 to 172.23.0.0/16.



Topic 62: Configure Floating Static Routes:

Floating static routes are static routes that have an administrative distance greater than the administrative distance of another static route or dynamic routes. They are very useful when providing a backup to a primary link.

By default, static routes have an administrative distance of 1, making them preferable to routes learned from dynamic routing protocols. For example, the administrative distances of some common dynamic routing protocols are:

EIGRP = 90

IGRP = 100

OSPF = 110

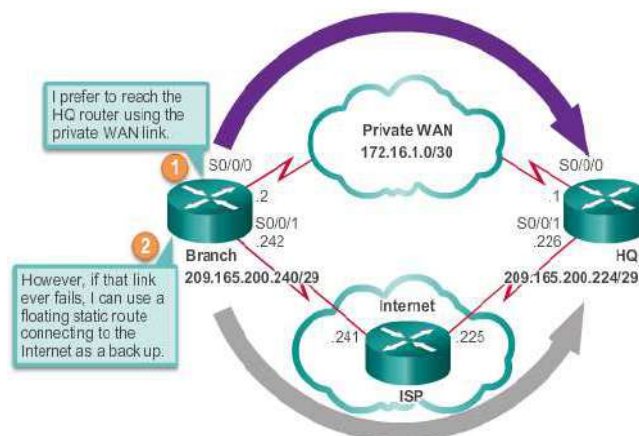
IS-IS = 115

RIP = 120

The administrative distance of a static route can be increased to make the route less desirable than that of another static route or a route learned through a dynamic routing protocol. In this way, the static route floats and is not used when the route with the better administrative distance is active. However, if the preferred route is lost, the floating static route can take over, and traffic can be sent through this alternate route.

A floating static route can be used to provide a backup route to multiple interfaces or networks on a router. It is also encapsulation independent, meaning it can be used to forward packets out any interface, regardless of encapsulation type.

An important consideration of a floating static route is that it is affected by convergence time. A route that is continuously dropping and re-establishing a connection can cause the backup interface to be activated unnecessarily.



Configuring a Floating Static Route

IPv4 static routes are configured using the **ip route** global configuration command and specifying an administrative distance. If no administrative distance is configured, the default value (1) is used.

See the slides to find out the example of configuring a floating static route.

Topic 63: Packet Tracer – How to Configure Static & Default Routes-1:

- Setup the Topology

- Configure Device Settings
- Configure Static Routes
- Configure Default Routes

Background/Scenario

In this activity, you will configure static and default routes. A static route is a route that is entered manually by the network administrator to create a route that is reliable and safe. There are four different static routes that are used in this activity: a recursive static route, a directly connected static route, a fully specified static route, and a default route.

Topic 64: Packet Tracer – How to Configure Static & Default Routes-2:

This is the continuation of the previous topic.

Topic 65: Packet Tracer – Configuring IPv4 Static & Default Routes:

- Configure Static and Default Routes
- Verify Connectivity

Background/Scenario

In this activity, you will configure static and default routes. A static route is a route that is entered manually by the network administrator to create a route that is reliable and safe. There are four different static routes that are used in this activity: a recursive static route, a directly connected static route, a fully specified static route, and a default route.

Topic 66: Packet Tracer – Configuring IPv4 Route Summarization:

- Calculate Summary Routes
- Configure Summary Routes
- Verify Connectivity

Topic 67: Packet Tracer – Configuring IPv4 Route Summarization - 2:

This topic is in continuation of the previous topic.

Topic 68: Packet Tracer – Configuring a Floating Static Route:

- Configure
- Test Failover to the Backup Route

Topic 69: Dynamic Routing Protocol:

Dynamic routing protocols have been used in networks since the late 1980s. One of the first routing protocols was Routing Information Protocol (RIP). RIP version 1 (RIPv1) was released in 1988. As networks evolved and became more complex, new routing protocols emerged. The RIP routing protocol was updated to accommodate growth in the network environment, into RIPv2. However, the newer version of RIP still does not scale to the larger network implementations of today. To address the needs

of larger networks, two advanced routing protocols were developed: Open Shortest Path First (OSPF) and Intermediate System-to-Intermediate System (IS-IS). Cisco developed the Interior Gateway Routing Protocol (IGRP) and Enhanced IGRP (EIGRP), which also scales well in larger network implementations.

Additionally, there was the need to connect different internetworks and provide routing between them. The Border Gateway Protocol (BGP) is now used between Internet service providers (ISPs). BGP is also used between ISPs and their larger private clients to exchange routing information.

	Interior Gateway Protocols				Exterior Gateway Protocols
	Distance Vector		Link-State		Path Vector
IPv4	RIPv2	EIGRP	OSPFv2	IS-IS	BGP-4
IPv6	RIPng	EIGRP for IPv6	OSPFv3	IS-IS for IPv6	BGP-MP

Routing protocols are used to facilitate the exchange of routing information between routers. A routing protocol is a set of processes, algorithms, and messages that are used to exchange routing information and populate the routing table with the routing protocols choice of best paths. The purpose of dynamic routing protocols includes:

- Discovery of remote networks
- Maintaining up-to-date routing information
- Choosing the best path to destination networks
- Ability to find a new best path if the current path is no longer available

The main components of dynamic routing protocols include:

Data structures - Routing protocols typically use tables or databases for its operations. This information is kept in RAM.

Routing protocol messages - Routing protocols use various types of messages to discover neighboring routers, exchange routing information, and other tasks to learn and maintain accurate information about the network.

Algorithm - An algorithm is a finite list of steps used to accomplish a task. Routing protocols use algorithms for facilitating routing information and for best path determination.

Topic 70: Routing Protocol Operating Fundamentals:

All routing protocols are designed to learn about remote networks and to quickly adapt whenever there is a change in the topology. The method that a routing protocol uses to accomplish this depends upon the algorithm it uses and the operational characteristics of that protocol.

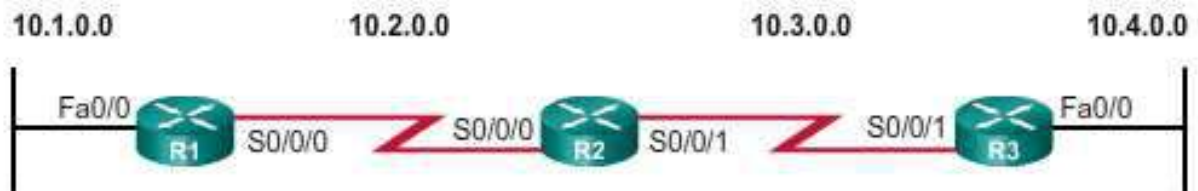
In general, the operations of a dynamic routing protocol can be described as follows:

1. The router sends and receives routing messages on its interfaces.
2. The router shares routing messages and routing information with other routers that are using the same routing protocol.

3. Routers exchange routing information to learn about remote networks.

4. When a router detects a topology change the routing protocol can advertise this change to other routers.

All routing protocols follow the same patterns of operation. To help illustrate this, consider the following scenario in which all three routers are running RIPv2.



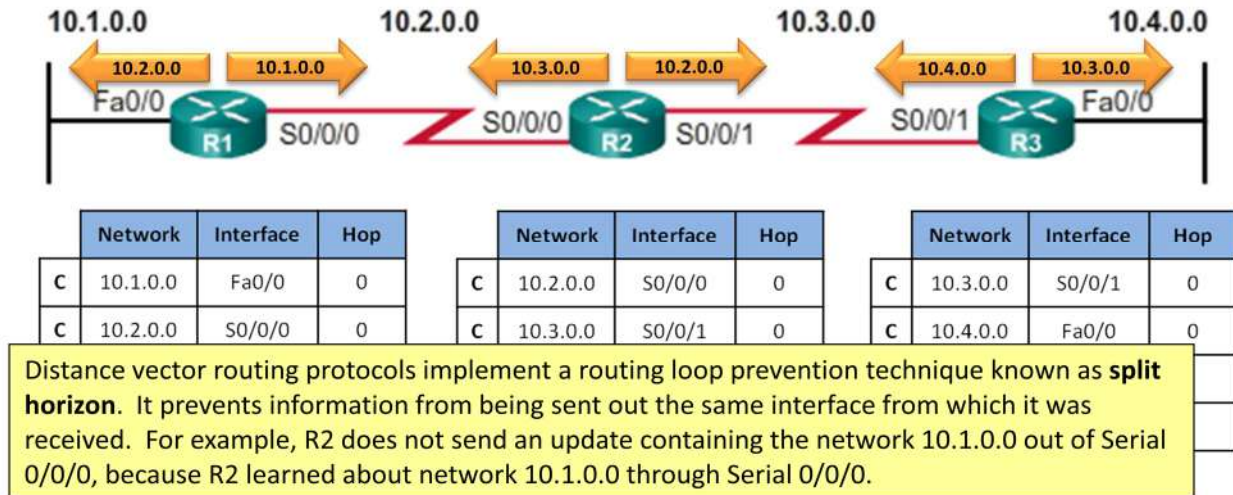
When a router powers up, it knows nothing about the network topology. It does not even know that there are devices on the other end of its links. The only information that a router has is from its own saved configuration file stored in NVRAM. After a router boots successfully, it applies the saved configuration. If the IP addressing is configured correctly, then the router initially discovers its own directly connected networks.

Routers proceed through the boot up process and then discovers any directly connected networks and subnet masks. This information is added to their routing tables as follows:

- R1 adds the 10.1.0.0 network available through interface FastEthernet 0/0 and 10.2.0.0 is available through interface Serial 0/0/0.
- R2 adds the 10.2.0.0 network available through interface Serial 0/0/0 and 10.3.0.0 is available through interface Serial 0/0/1.
- R3 adds the 10.3.0.0 network available through interface Serial 0/0/1 and 10.4.0.0 is available through interface FastEthernet 0/0.

With this initial information, the routers then proceed to find additional route sources for their routing tables.

After initial boot up and discovery, the routing table is updated with all directly connected networks and the interfaces those networks reside on.



If a routing protocol is configured, the next step is for the router to begin exchanging routing updates to learn about any remote routes.

The router sends an update packet out all interfaces that are enabled on the router. The update contains the information in the routing table, which currently are all directly connected networks.

At the same time, the router also receives and processes similar updates from other connected routers. Upon receiving an update, the router checks it for new network information. Any networks that are not currently listed in the routing table are added.

Refer to the figure above for a topology setup between three routers, R1, R2, and R3. Based on this topology, below is a listing of the different updates that R1, R2, and R3 send and receive during initial convergence.

R1:

- Sends an update about network 10.1.0.0 out the Serial0/0/0 interface
- Sends an update about network 10.2.0.0 out the FastEthernet0/0 interface
- Receives update from R2 about network 10.3.0.0 and increments the hop count by 1
- Stores network 10.3.0.0 in the routing table with a metric of 1

R2:

- Sends an update about network 10.3.0.0 out the Serial 0/0/0 interface
- Sends an update about network 10.2.0.0 out the Serial 0/0/1 interface
- Receives an update from R1 about network 10.1.0.0 and increments the hop count by 1
- Stores network 10.1.0.0 in the routing table with a metric of 1
- Receives an update from R3 about network 10.4.0.0 and increments the hop count by 1
- Stores network 10.4.0.0 in the routing table with a metric of 1

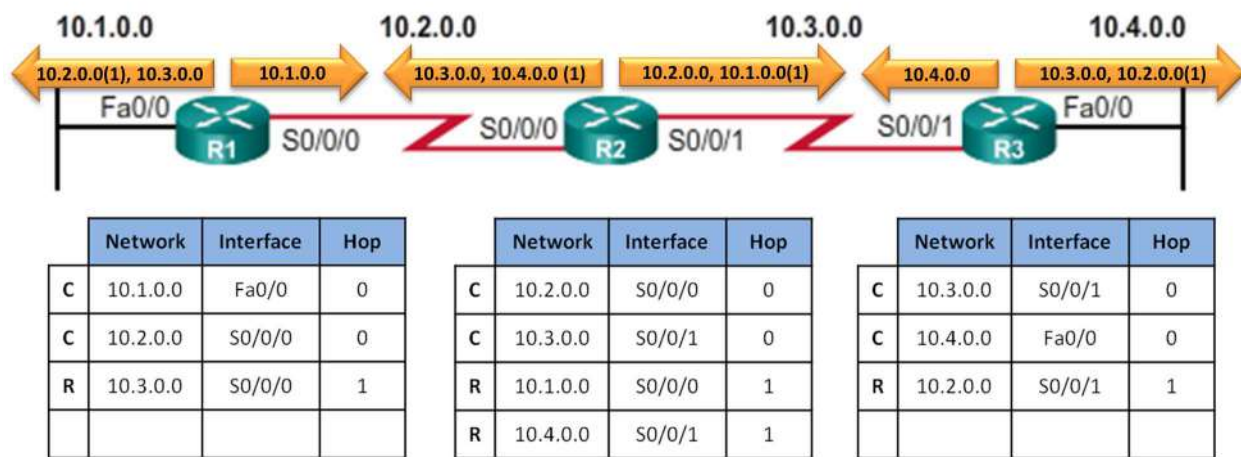
R3:

- Sends an update about network 10.4.0.0 out the Serial 0/0/1 interface

- Sends an update about network 10.3.0.0 out the FastEthernet0/0
- Receives an update from R2 about network 10.2.0.0 and increments the hop count by 1
- Stores network 10.2.0.0 in the routing table with a metric of 1

After this first round of update exchanges, each router knows about the connected networks of their directly connected neighbors. However, did you notice that R1 does not yet know about 10.4.0.0 and that R3 does not yet know about 10.1.0.0? Full knowledge and a converged network do not take place until there is another exchange of routing information.

At this point the routers have knowledge about their own directly connected networks and about the connected networks of their immediate neighbors. Continuing the journey toward convergence, the routers exchange the next round of periodic updates. Each router again checks the updates for new information.



Refer to the figure above for a topology setup between three routers, R1, R2, and R3. After initial discovery is complete, each router continues the convergence process by sending and receiving the following updates.

R1:

- Sends an update about network 10.1.0.0 out the Serial 0/0/0 interface
- Sends an update about networks 10.2.0.0 and 10.3.0.0 out the FastEthernet0/0 interface
- Receives an update from R2 about network 10.4.0.0 and increments the hop count by 1
- Stores network 10.4.0.0 in the routing table with a metric of 2
- Same update from R2 contains information about network 10.3.0.0 with a metric of 1. There is no change; therefore, the routing information remains the same

R2:

- Sends an update about networks 10.3.0.0 and 10.4.0.0 out of Serial 0/0/0 interface
- Sends an update about networks 10.1.0.0 and 10.2.0.0 out of Serial 0/0/1 interface
- Receives an update from R1 about network 10.1.0.0. There is no change; therefore, the routing information remains the same

- Receives an update from R3 about network 10.4.0.0. There is no change; therefore, the routing information remains the same

R3:

- Sends an update about network 10.4.0.0 out the Serial 0/0/1 interface
- Sends an update about networks 10.2.0.0 and 10.3.0.0 out the FastEthernet0/0 interface
- Receives an update from R2 about network 10.1.0.0 and increments the hop count by 1
- Stores network 10.1.0.0 in the routing table with a metric of 2
- Same update from R2 contains information about network 10.2.0.0 with a metric of 1. There is no change; therefore, the routing information remains the same

Distance vector routing protocols typically implement a routing loop prevention technique known as split horizon. Split horizon prevents information from being sent out the same interface from which it was received. For example, R2 does not send an update containing the network 10.1.0.0 out of Serial 0/0/0, because R2 learned about network 10.1.0.0 through Serial 0/0/0.

After routers within a network have converged, the router can then use the information within the route table to determine the best path to reach a destination. Different routing protocols have different ways of calculating the best path.

The network has converged when all routers have complete and accurate information about the entire network. Convergence time is the time it takes routers to share information, calculate best paths, and update their routing tables. A network is not completely operable until the network has converged; therefore, most networks require short convergence times.

Convergence is both collaborative and independent. The routers share information with each other, but must independently calculate the impacts of the topology change on their own routes. Because they develop an agreement with the new topology independently, they are said to converge on this consensus.

Convergence properties include the speed of propagation of routing information and the calculation of optimal paths. The speed of propagation refers to the amount of time it takes for routers within the network to forward routing information.

Routing protocols can be rated based on the speed to convergence; the faster the convergence, the better the routing protocol. Generally, older protocols, such as RIP, are slow to converge, whereas modern protocols, such as EIGRP and OSPF, converge more quickly.

Topic 71: Packet Tracer – Investigating Convergence:

- View the Routing Table of a Converged Network
- Add a New LAN
- Watch Network Converge

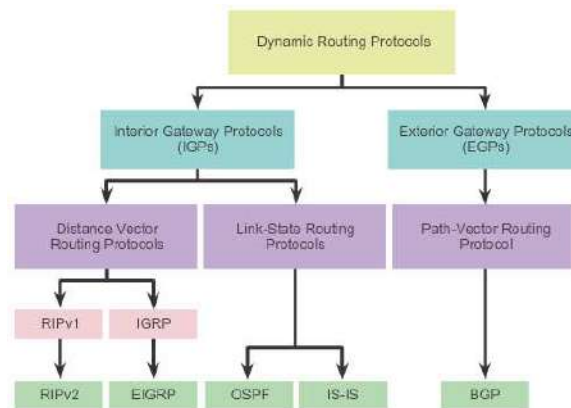
Topic 72: Types of Routing Protocols:

Routing protocols can be classified into different groups according to their characteristics. Specifically, routing protocols can be classified by their:

- **Purpose** - Interior Gateway Protocol (IGP) or Exterior Gateway Protocol (EGP)
- **Operation** - Distance vector, link-state protocol, or path-vector protocol
- **Behavior** - Classful (legacy) or classless protocol
- For example, IPv4 routing protocols are classified as follows:
- **RIPv1 (legacy)** - IGP, distance vector, classful protocol
- **IGRP (legacy)** - IGP, distance vector, classful protocol developed by Cisco (deprecated from 12.2 IOS and later)
- **RIPv2** - IGP, distance vector, classless protocol
- **EIGRP** - IGP, distance vector, classless protocol developed by Cisco
- **OSPF** - IGP, link-state, classless protocol
- **IS-IS** - IGP, link-state, classless protocol
- **BGP** - EGP, path-vector, classless protocol

The classful routing protocols, RIPv1 and IGRP, are legacy protocols and are only used in older networks. These routing protocols have evolved into the classless routing protocols, RIPv2 and EIGRP, respectively. Link-state routing protocols are classless by nature.

Figure below displays a hierarchical view of dynamic routing protocol classification.



An autonomous system (AS) is a collection of routers under a common administration such as a company or an organization. An AS is also known as a routing domain. Typical examples of an AS are a company's internal network and an ISP's network.

The Internet is based on the AS concept; therefore, two types of routing protocols are required:

Interior Gateway Protocols (IGP) - Used for routing within an AS. It is also referred to as intra-AS routing. Companies, organizations, and even service providers use an IGP on their internal networks. IGPs include RIP, EIGRP, OSPF, and IS-IS.

Exterior Gateway Protocols (EGP) - Used for routing between AS. It is also referred to as inter-AS routing. Service providers and large companies may interconnect using an EGP. The Border Gateway Protocol (BGP) is the only currently-viable EGP and is the official routing protocol used by the Internet.

Note: Because BGP is the only EGP available, the term EGP is rarely used; instead, most engineers simply refer to BGP.

Distance Vector Routing Protocols

Distance vector means that routes are advertised by providing two characteristics:

Distance - Identifies how far it is to the destination network and is based on a metric such as the hop count, cost, bandwidth, delay, and more.

Vector - Specifies the direction of the next-hop router or exit interface to reach the destination.

A router using a distance vector routing protocol does not have the knowledge of the entire path to a destination network. Distance vector protocols use routers as sign posts along the path to the final destination. The only information a router knows about a remote network is the distance or metric to reach that network and which path or interface to use to get there. Distance vector routing protocols do not have an actual map of the network topology.

There are four distance vector IPv4 IGPs:

- **RIPv1** - First generation legacy protocol
- **RIPv2** - Simple distance vector routing protocol
- **IGRP** - First generation Cisco proprietary protocol (obsolete and replaced by EIGRP)
- **EIGRP** - Advanced version of distance vector routing

Link State Routing Protocols

In contrast to distance vector routing protocol operation, a router configured with a link-state routing protocol can create a complete view or topology of the network by gathering information from all of the other routers.

RIP-enabled routers send periodic updates of their routing information to their neighbors. Link-state routing protocols do not use periodic updates. After the network has converged, a link-state update is only sent when there is a change in the topology.

Link-state protocols work best in situations where:

- The network design is hierarchical, usually occurring in large networks
- Fast convergence of the network is crucial
- The administrators have good knowledge of the implemented link-state routing protocol

There are two link-state IPv4 IGPs:

- **OSPF** - Popular standards based routing protocol
- **IS-IS** - Popular in provider networks

Topic 73: Classful and Classless Routing Protocols:

- Classful routing protocols do not send subnet mask information in their routing updates:
 - Only RIPv1 and IGRP are classful.
 - Created when network addresses were allocated based on classes (class A, B, or C).

- Cannot provide variable length subnet masks (VLSMs) and classless interdomain routing (CIDR).
- Create problems in discontinuous networks.

Topic 74: Distance Vector Dynamic Routing:

Distance vector routing protocols share updates between neighbors. Neighbors are routers that share a link and are configured to use the same routing protocol. The router is only aware of the network addresses of its own interfaces and the remote network addresses it can reach through its neighbors. Routers using distance vector routing are not aware of the network topology.

Some distance vector routing protocols send periodic updates. For example, RIP sends a periodic update to all of its neighbors every 30 seconds. RIP does this even if the topology has not changed; it continues to send updates. RIPv1 reaches all of its neighbors by sending updates to the all-hosts IPv4 address of 255.255.255.255, a broadcast.

The broadcasting of periodic updates is inefficient because the updates consume bandwidth and consume network device CPU resources. Every network device has to process a broadcast message. RIPv2 and EIGRP, instead, use multicast addresses so that only neighbors that need updates will receive them. EIGRP can also send a unicast message to only the affected neighbor. Additionally, EIGRP only sends an update when needed, instead of periodically.

The two modern IPv4 distance vector routing protocols are RIPv2 and EIGRP. RIPv1 and IGRP are listed only for historical accuracy.

IGRP

The Interior Gateway Routing Protocol (IGRP) was the first proprietary IPv4 routing protocol developed by Cisco in 1984. It used the following design characteristics:

- Bandwidth, delay, load, and reliability are used to create a composite metric.
- Routing updates are broadcast every 90 seconds, by default.

In 1992, IGRP was replaced by Enhanced IGRP (EIGRP). Like RIPv2, EIGRP also introduced support for VLSM and CIDR. EIGRP increases efficiency, reduces routing updates, and supports secure message exchange.

EIGRP also introduced:

Bounded triggered updates - It does not send periodic updates. Only routing table changes are propagated, whenever a change occurs. This reduces the amount of load the routing protocol places on the network. Bounded triggered updates means that EIGRP only sends to the neighbors that need it. It uses less bandwidth, especially in large networks with many routes.

Hello keepalive mechanism - A small Hello message is periodically exchanged to maintain adjacencies with neighboring routers. This means a very low usage of network resources during normal operation, instead of the periodic updates.

Maintains a topology table - Maintains all the routes received from neighbors (not only the best paths) in a topology table. DUAL can insert backup routes into the EIGRP topology table.

Rapid convergence - In most cases, it is the fastest IGP to converge because it maintains alternate routes, enabling almost instantaneous convergence. If a primary route fails, the router can use the alternate route identified. The switchover to the alternate route is immediate and does not involve interaction with other routers.

Topic 75: RIP Protocol - 1:

Although RIP is rarely used in modern networks, it is useful as a foundation for understanding basic network routing. For this reason, this section provides a brief overview of how to configure basic RIP settings and to verify RIPv2.

Please see the slides for more information about RIP configuration and other parameters.

Topic 76: RIP Protocol - 2:

RIP Versions

By default, when a RIP process is configured on a Cisco router, it is running RIPv1. However, even though the router only sends RIPv1 messages, it can interpret both RIPv1 and RIPv2 messages. A RIPv1 router ignores the RIPv2 fields in the route entry.

Use the **version 2** router configuration mode command to enable RIPv2. Please see the slides for information about how to use RIP version 2.

Disabling Auto Summarization

RIPv2 automatically summarizes networks at major network boundaries by default, just like RIPv1.

To modify the default RIPv2 behavior of automatic summarization, use the **no auto-summary** router configuration mode command. This command has no effect when using RIPv1. When automatic summarization has been disabled, RIPv2 no longer summarizes networks to their classful address at boundary routers. RIPv2 now includes all subnets and their appropriate masks in its routing updates. The **show ip protocols** now states that automatic network summarization is not in effect.

Please see the slides for example of auto-summarization.

RIP Passive Interfaces

By default, RIP updates are forwarded out all RIP enabled interfaces. However, RIP updates really only need to be sent out interfaces connecting to other RIP enabled routers.

Sending out unneeded updates on a LAN impacts the network in three ways:

Wasted Bandwidth - Bandwidth is used to transport unnecessary updates. Because RIP updates are either broadcasted or multicasted; therefore, switches also forward the updates out all ports.

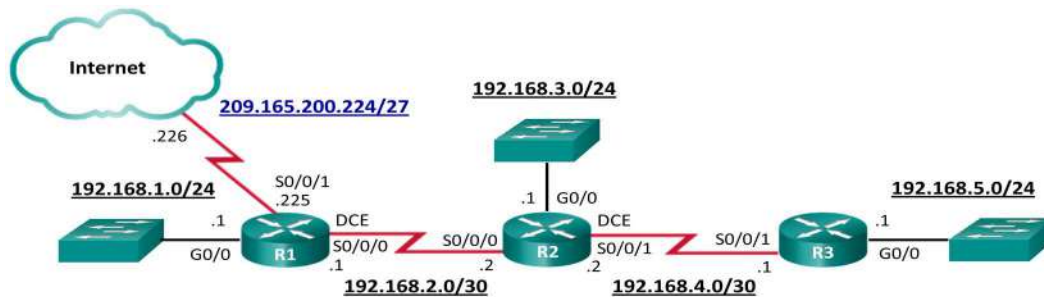
Wasted Resources - All devices on the LAN must process the update up to the transport layers, at which point the devices will discard the update.

Security Risk - Advertising updates on a broadcast network is a security risk. RIP updates can be intercepted with packet sniffing software. Routing updates can be modified and sent back to the router, corrupting the routing table with false metrics that misdirect traffic.

Use the **passive-interface** router configuration command to prevent the transmission of routing updates through a router interface, but still allow that network to be advertised to other routers. The command stops routing updates out the specified interface. However, the network that the specified interface belongs to is still advertised in routing updates that are sent out other interfaces.

Please see the slides for an example of passive interfaces.

Propagating a Default Route



Refer to Figure above. In this scenario, R1 is single-homed to a service provider. Therefore, all that is required for R1 to reach the Internet is a default static route going out of the Serial 0/0/1 interface.

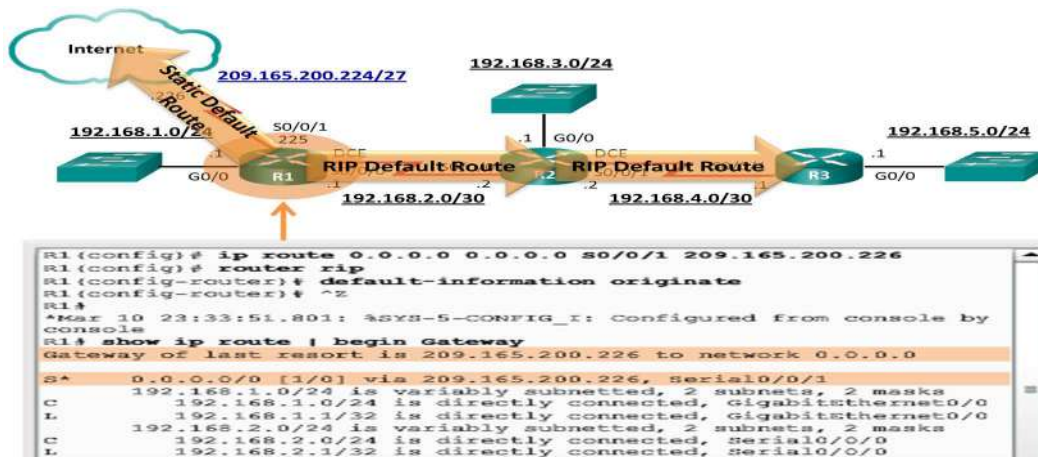
Similar default static routes could be configured on R2 and R3, but it is much more scalable to enter it one time on the edge router R1 and then have R1 propagate it to all other routers using RIP. To provide Internet connectivity to all other networks in the RIP routing domain, the default static route needs to be advertised to all other routers that use the dynamic routing protocol.

To propagate a default route, the edge router must be configured with:

A default static route using the **ip route 0.0.0.0 0.0.0.0 exit-intf next-hop-ip** command.

The **default-information originate** router configuration command. This instructs R1 router to originate default information, by propagating the static default route in RIP updates.

The example in Figure below configures a fully specified default static route to the service provider and then the route is propagated by RIP. Notice that R1 now has a Gateway of Last Resort and default route installed in its routing table.



Topic 77: Packet Tracer – Configuring RIPv2-1:

- Build the Network
- Configure Device Settings
- Configure RIPv2
- Verify RIPv2

Background/Scenario

In this activity, you will configure static and default routes. A static route is a route that is entered manually by the network administrator to create a route that is reliable and safe. There are four different static routes that are used in this activity: a recursive static route, a directly connected static route, a fully specified static route, and a default route.

Topic 78: Packet Tracer – Configuring RIPv2-2:

This is the continuation of the previous topic.

Topic 79: Packet Tracer – Configuring RIPv2-3:

This is the continuation of the previous topic.

Topic 80: Packet Tracer – Comparing RIP and EIGRP Path Selection:

- Predict the Path
- Trace the Route

Topic 81: Packet Tracer – Configuring RIPv2:

- Configure RIPv2
- Verify Configurations

Topic 82: Link-State Routing Protocol:

Link-state routing protocols are also known as shortest path first protocols and are built around Edsger Dijkstra's shortest path first (SPF) algorithm.

Link-state routing protocols have the reputation of being much more complex than their distance vector counterparts. However, the basic functionality and configuration of link-state routing protocols is equally straight-forward.

Just like RIP and EIGRP, basic OSPF operations can be configured using the:

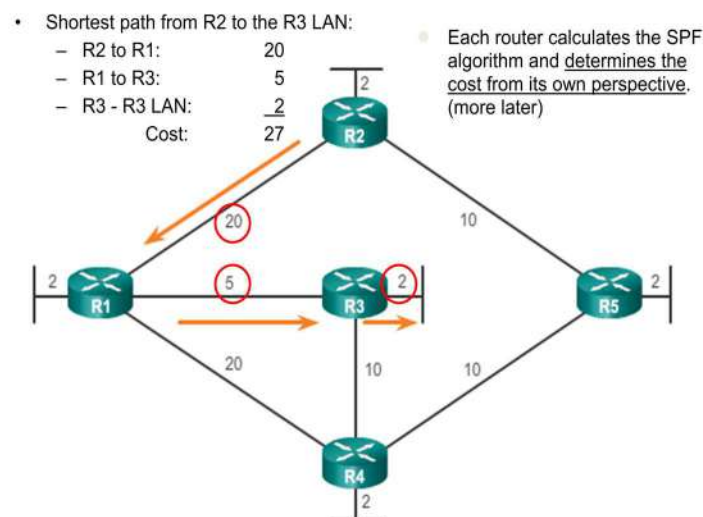
router ospf process-id global configuration command

network command to advertise networks

Dijkstra's Shortest Path First Algorithm

All link-state routing protocols apply Dijkstra's algorithm to calculate the best path route. The algorithm is commonly referred to as the shortest path first (SPF) algorithm. This algorithm uses accumulated costs along each path, from source to destination, to determine the total cost of a route.

In the figure below, each path is labeled with an arbitrary value for cost. The cost of the shortest path for R2 to send packets to the LAN attached to R3 is 27. Each router determines its own cost to each destination in the topology. In other words, each router calculates the SPF algorithm and determines the cost from its own perspective.



Topic 83: Link State Updates - 1:

So exactly how does a link-state routing protocol work? With link-state routing protocols, a link is an interface on a router. Information about the state of those links is known as link-states.

Examine the topology in the figure below. All routers in the topology will complete the following generic link-state routing process to reach a state of convergence:

1. Each router learns about its own links and its own directly connected networks. This is done by detecting that an interface is in the up state.

2. Each router is responsible for meeting its neighbors on directly connected networks. Link state routers do this by exchanging Hello packets with other link-state routers on directly connected networks.

3. Each router builds a Link-State Packet (LSP) containing the state of each directly connected link. This is done by recording all the pertinent information about each neighbor, including neighbor ID, link type, and bandwidth.

4. Each router floods the LSP to all neighbors. Those neighbors store all LSPs received in a database. They then flood the LSPs to their neighbors until all routers in the area have received the LSPs. Each router stores a copy of each LSP received from its neighbors in a local database.

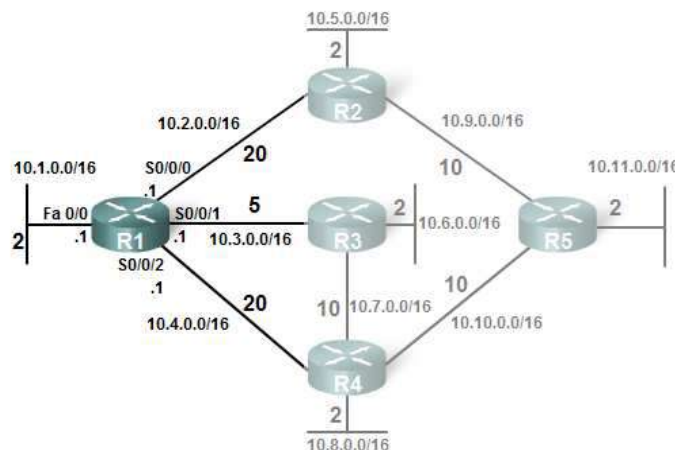
5. Each router uses the database to construct a complete map of the topology and computes the best path to each destination network. Like having a road map, the router now has a complete map of all destinations in the topology and the routes to reach them. The SPF algorithm is used to construct the map of the topology and to determine the best path to each network.

Note: This process is the same for both OSPF for IPv4 and OSPF for IPv6. The examples in this section refer to OSPF for IPv4.

Step 1: Learning About Directly Connected Networks

The first step in the link-state routing process is that each router learns about its own links, its own directly connected networks. When a router interface is configured with an IP address and subnet mask, the interface becomes part of that network.

As with distance vector protocols and static routes, the interface must be properly configured with an IPv4 address and subnet mask, and the link must be in the up state before the link-state routing protocol can learn about a link. Also, like distance vector protocols, the interface must be included in one of the **network** router configuration statements before it can participate in the link-state routing process.



Step 2: Sending Hello Packets to Neighbors

The second step in the link-state routing process is that each router is responsible for meeting its neighbors on directly connected networks.

Routers with link-state routing protocols use a Hello protocol to discover any neighbors on its links. A neighbor is any other router that is enabled with the same link-state routing protocol.

When two link-state routers learn that they are neighbors, they form an adjacency. These small Hello packets continue to be exchanged between two adjacent neighbors and serves as a keepalive function to monitor the state of the neighbor. If a router stops receiving Hello packets from a neighbor, that neighbor is considered unreachable and the adjacency is broken.

Please see the slides for working example of Link state routing protocol.

Topic 84: Link State Updates - 2:

Step 3: Building the Link State Packet

The third step in the link-state routing process is that each router builds a link-state packet (LSP) containing the state of each directly connected link.

After a router has established its adjacencies, it can build its LSPs that contain the link-state information about its links.

Step 4: Flooding the LSP

The fourth step in the link-state routing process is that each router floods the LSP to all neighbors, who then store all LSPs received in a database.

Each router floods its link-state information to all other link-state routers in the routing area. Whenever a router receives an LSP from a neighboring router, it immediately sends that LSP out all other interfaces except the interface that received the LSP. This process creates a flooding effect of LSPs from all routers throughout the routing area.

Remember that LSPs do not need to be sent periodically. An LSP only needs to be sent:

- During initial startup of the routing protocol process on that router (e.g., router restart)
- Whenever there is a change in the topology (e.g., a link going down or coming up, a neighbor adjacency being established or broken)

In addition to the link-state information, other information is included in the LSP, such as sequence numbers and aging information, to help manage the flooding process. This information is used by each router to determine if it has already received the LSP from another router or if the LSP has newer information than what is already contained in the link-state database. This process allows a router to keep only the most current information in its link-state database.

Step 5: Building the Link State Database

The final step in the link-state routing process is that each router uses the database to construct a complete map of the topology and computes the best path to each destination network.

Eventually, all routers receive an LSP from every other link-state router in the routing area. These LSPs are stored in the link-state database.

Building the SPF Tree

Each router in the routing area uses the link-state database and SPF algorithm to construct the SPF tree.

Each router constructs its own SPF tree independently from all other routers. To ensure proper routing, the link-state databases used to construct those trees must be identical on all routers.

Step 5: Generating Routing Table from SPF Tree

Using the shortest path information determined by the SPF algorithm, these paths can now be added to the routing table. The routing table also includes all directly connected networks and routes from any other sources, such as static routes. Packets are now forwarded according to these entries in the routing table.

Please see the slides for working example of Link state routing protocol.

Topic 85: Why Use Link State Routing Protocol:

Advantages of Link State Routing Protocols

There are several advantages of link-state routing protocols compared to distance vector routing protocols.

Builds a Topological Map - Link-state routing protocols create a topological map, or SPF tree of the network topology. Because link-state routing protocols exchange link-states, the SPF algorithm can build an SPF tree of the network. Using the SPF tree, each router can independently determine the shortest path to every network.

Fast Convergence - When receiving an LSP, link-state routing protocols immediately flood the LSP out all interfaces except for the interface from which the LSP was received. In contrast, RIP needs to process each routing update and update its routing table before flooding them out other interfaces.

Event-driven Updates - After the initial flooding of LSPs, link-state routing protocols only send out an LSP when there is a change in the topology. The LSP contains only the information regarding the affected link. Unlike some distance vector routing protocols, link-state routing protocols do not send periodic updates.

Hierarchical Design - Link-state routing protocols use the concept of areas. Multiple areas create a hierarchical design to networks, allowing for better route aggregation (summarization) and the isolation of routing issues within an area.

Link-state protocols also have a few disadvantages compared to distance vector routing protocols:

Memory Requirements - Link-state protocols require additional memory to create and maintain the link-state database and SPF tree.

Processing Requirements - Link-state protocols can also require more CPU processing than distance vector routing protocols. The SPF algorithm requires more CPU time than distance vector algorithms such as Bellman-Ford, because link-state protocols build a complete map of the topology.

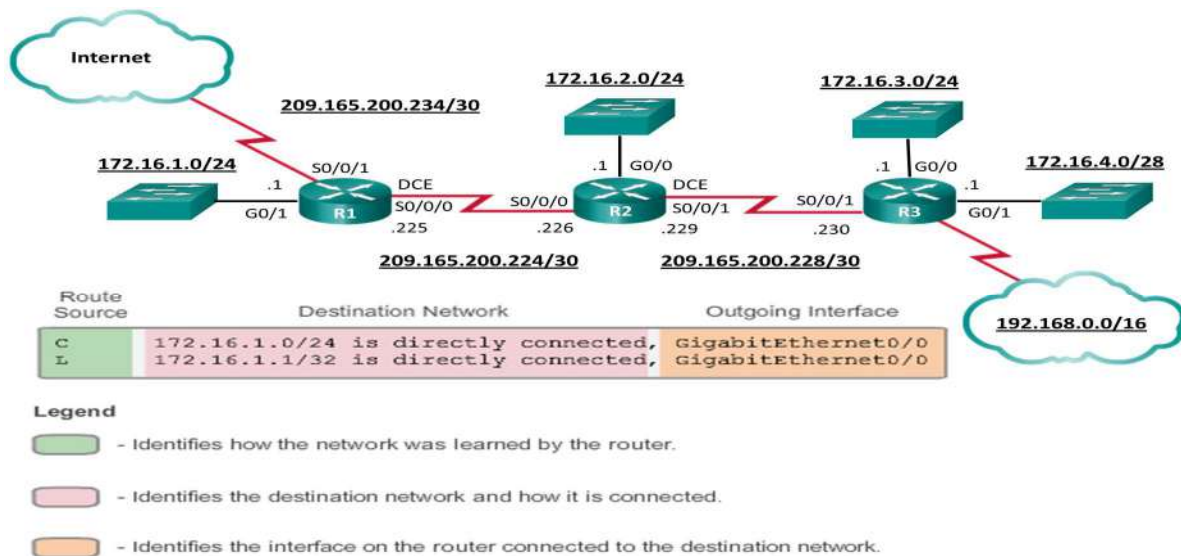
Bandwidth Requirements - The flooding of link-state packets can adversely affect the available bandwidth on a network. This should only occur during initial startup of routers, but can also be an issue on unstable networks.

Addressing Disadvantages

Modern link-state routing protocols are designed to minimize the effects on memory, CPU, and bandwidth. The use and configuration of multiple areas can reduce the size of the link-state databases. Multiple areas can also limit the amount of link-state information flooding in a routing domain and send LSPs only to those routers that need them. When there is a change in the topology, only those routers in the affected area receive the LSP and run the SPF algorithm. This can help isolate an unstable link to a specific area in the routing domain.

Topic 86: The Routing Table:

Reference Topology – Connected and Local



The topology displayed in Figure above is used as the reference topology. Notice that in the topology:

R1 is the edge router that connects to the Internet. Therefore, it is propagating a default static route to R2 and R3.

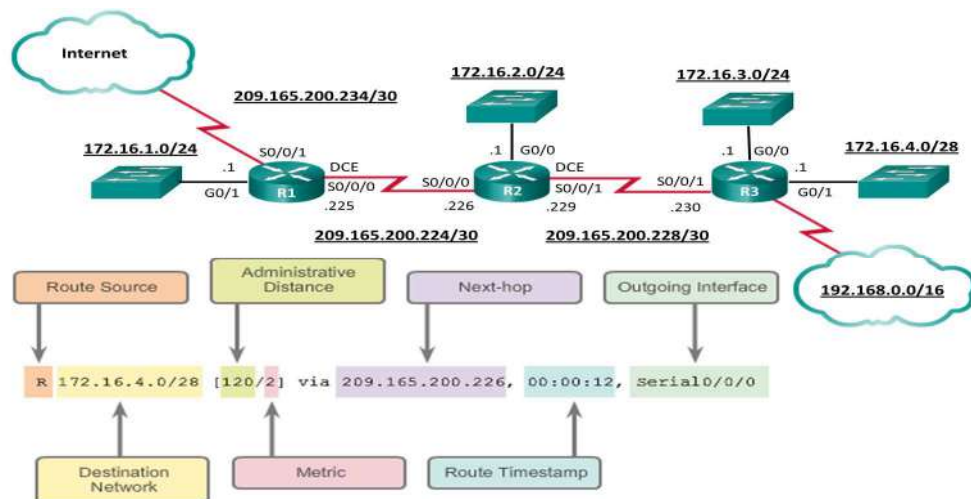
R1, R2, and R3 contain discontinuous networks separated by another classful network.

R3 is also introducing a 192.168.0.0/16 supernet route.

Figure above displays the IPv4 routing table of R1 with directly connected, static, and dynamic routes.

Note: The routing table hierarchy in Cisco IOS was originally implemented with the classful routing scheme. Although the routing table incorporates both classful and classless addressing, the overall structure is still built around this classful scheme.

Remote Networks



The figure above displays an IPv4 routing table entry on R1 for the route to remote network 172.16.4.0 on R3. The entry identifies the following information:

- **Route source** - Identifies how the route was learned.
- **Destination network** - Identifies the address of the remote network.
- **Administrative distance** - Identifies the trustworthiness of the route source.
- **Metric** - Identifies the value assigned to reach the remote network. Lower values indicate preferred routes.
- **Next hop** - Identifies the IPv4 address of the next router to forward the packet to.
- **Route timestamp** - Identifies from when the route was last heard.
- **Outgoing interface** - Identifies the exit interface to use to forward a packet toward the final destination.

IPv4 Routing Table – Ultimate Route

A dynamically built routing table provides a great deal of information. Therefore, it is crucial to understand the output generated by the routing table. Special terms are applied when discussing the contents of a routing table.

The Cisco IP routing table is not a flat database. The routing table is actually a hierarchical structure that is used to speed up the lookup process when locating routes and forwarding packets. Within this structure, the hierarchy includes several levels.

Routes are discussed in terms of:

- Ultimate route
- Level 1 route
- Level 1 parent route
- Level 2 child routes

IPv4 Routing Table – Level 1 Route

A level 1 route is a route with a subnet mask equal to or less than the classful mask of the network address. Therefore, a level 1 route can be a:

Network route - A network route that has a subnet mask equal to that of the classful mask.

Supernet route - A supernet route is a network address with a mask less than the classful mask, for example, a summary address.

Default route - A default route is a static route with the address 0.0.0.0/0.

The source of the level 1 route can be a directly connected network, static route, or a dynamic routing protocol.

IPv4 Routing Table – Level 2 Child Route

A level 2 child route is a route that is a subnet of a classful network address. A level 1 parent route is a level 1 network route that is subnetted. A level 1 parent routes contain level 2 child routes.

Like a level 1 route, the source of a level 2 route can be a directly connected network, a static route, or a dynamically learned route. Level 2 child routes are also ultimate routes.

Please see the slides for an example of it.

Topic 87: IPv4 Route Lookup Process:

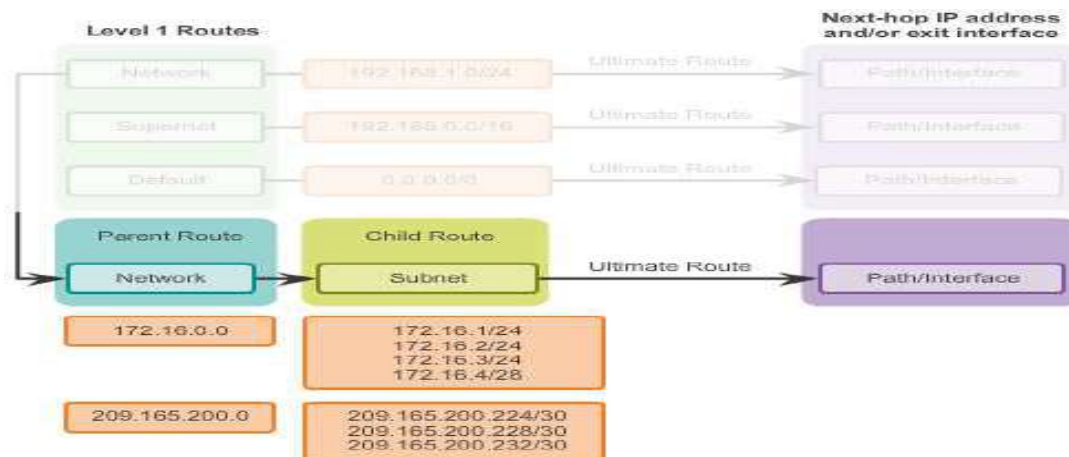
When a packet arrives on a router interface, the router examines the IPv4 header, identifies the destination IPv4 address, and proceeds through the router lookup process.

In Figure below, the router examines level 1 network routes for the best match with the destination address of the IPv4 packet.

1. If the best match is a level 1 ultimate route, then this route is used to forward the packet.



2. If the best match is a level 1 parent route, proceed to the next step.



In Figure above, the router examines child routes (the subnet routes) of the parent route for a best match.

3. If there is a match with a level 2 child route, that subnet is used to forward the packet.



4. If there is not a match with any of the level 2 child routes, proceed to the next step.

In Figure above, the router continues searching level 1 supernet routes in the routing table for a match, including the default route, if there is one.

5. If there is now a lesser match with a level 1 supernet or default routes, the router uses that route to forward the packet.

6. If there is not a match with any route in the routing table, the router drops the packet.

Note: A route referencing only a next-hop IP address and not an exit interface must be resolved to a route with an exit interface. A recursive lookup is performed on the next-hop IP address until the route is resolved to an exit interface.

Best Route = Longest Match

What is meant by the router must find the best match in the routing table? Best match is equal to the longest match.

For there to be a match between the destination IPv4 address of a packet and a route in the routing table, a minimum number of far left bits must match between the IPv4 address of the packet and the route in the routing table. The subnet mask of the route in the routing table is used to determine the minimum number of far left bits that must match. Remember that an IPv4 packet only contains the IPv4 address and not the subnet mask.

The best match is the route in the routing table that has the most number of far left matching bits with the destination IPv4 address of the packet. The route with the greatest number of equivalent far left bits, or the longest match, is always the preferred route.

In the figure below, a packet is destined for 172.16.0.10. The router has three possible routes that match this packet: 172.16.0.0/12, 172.16.0.0/18, and 172.16.0.0/26. Of the three routes, 172.16.0.0/26 has the longest match and is therefore chosen to forward the packet. Remember, for any of these routes to be considered a match there must be at least the number of matching bits indicated by the subnet mask of the route.

IP Packet Destination	172.16.0.10	10101100.00010000.00000000.00001010
Route 1	172.16.0.0/12	10101100.00010000.00000000.00000000
Route 2	172.16.0.0/18	10101100.00010000.00000000.00000000
Route 3	172.16.0.0/26	10101100.00010000.00000000.00000000

Longest Match to IP Packet Destination 

Topic 88: Open Shortest Path First:

	Interior Gateway Protocols		Exterior Gateway Protocols	
	Distance Vector Routing Protocols	Link State Routing Protocols	Path Vector	
Classful	RIP	IGRP		EGP
Classless	RIPv2	EIGRP	OSPFv2	IS-IS
IPv6	RIPng	EIGRP for IPv6	OSPFv3	IS-IS for IPv6
				BGPv4 for IPv6

As shown in Figure above, OSPF version 2 (OSPFv2) is available for IPv4 while OSPF version 3 (OSPFv3) is available for IPv6.

The initial development of OSPF began in 1987 by the Internet Engineering Task Force (IETF) OSPF Working Group. At that time, the Internet was largely an academic and research network funded by the U.S. government.

In 1989, the specification for OSPFv1 was published in RFC 1131. Two implementations were written. One implementation was developed to run on routers and the other to run on UNIX workstations. The

latter implementation became a widespread UNIX process known as GATED. OSPFv1 was an experimental routing protocol and was never deployed.

In 1991, OSPFv2 was introduced in RFC 1247 by John Moy. OSPFv2 offered significant technical improvements over OSPFv1. It is classless by design; therefore, it supports VLSM and CIDR.

At the same time the OSPF was introduced, ISO was working on a link-state routing protocol of their own, Intermediate System-to-Intermediate System (IS-IS). IETF chose OSPF as their recommended Interior Gateway Protocol (IGP).

In 1998, the OSPFv2 specification was updated in RFC 2328, which remains the current RFC for OSPF.

In 2008, OSPFv3 was updated in RFC 5340 as OSPF for IPv6.

Features of OSPF



OSPF features, as shown in Figure above, include:

Classless - It is classless by design; therefore, it supports VLSM and CIDR.

Efficient - Routing changes trigger routing updates (no periodic updates). It uses the SPF algorithm to choose the best path.

Fast convergence - It quickly propagates network changes.

Scalable - It works well in small and large network sizes. Routers can be grouped into areas to support a hierarchical system.

Secure - It supports Message Digest 5 (MD5) authentication. When enabled, OSPF routers only accept encrypted routing updates from peers with the same pre-shared password.

Administrative distance (AD) is the trustworthiness (or preference) of the route source. OSPF has a default administrative distance of 110.

Components of OSPF

All routing protocols share similar components. They all use routing protocol messages to exchange route information. The messages help build data structures, which are then processed using a routing algorithm.

The three main components of the OSPF routing protocol include:

Data Structures

OSPF creates and maintains three databases:

- **Adjacency database** - Creates the neighbor table
- **Link-state database (LSDB)** - Creates the topology table
- **Forwarding database** - Creates the routing table

These tables contain a list of neighboring routers to exchange routing information with and are kept and maintained in RAM.

Routing Protocol Messages

OSPF exchanges messages to convey routing information using five types of packets.

- Hello packet
- Database description packet
- Link-state request packet
- Link-state update packet
- Link-state acknowledgment packet

These packets are used to discover neighboring routers and also to exchange routing information to maintain accurate information about the network.

Algorithm

The CPU processes the neighbor and topology tables using Dijkstra's SPF algorithm. The SPF algorithm is based on the cumulative cost to reach a destination.

The SPF algorithm creates an SPF tree by placing each router at the root of the tree and calculating the shortest path to each node. The SPF tree is then used to calculate the best routes. OSPF places the best routes into the forwarding database, which is used to make the routing table.

Link State Operation

To maintain routing information, OSPF routers complete the following generic link-state routing process to reach a state of convergence:

1. **Establish Neighbor Adjacencies** - OSPF-enabled routers must recognize each other on the network before they can share information. An OSPF-enabled router sends Hello packets out all OSPF-enabled interfaces to determine if neighbors are present on those links. If a neighbor is present, the OSPF-enabled router attempts to establish a neighbor adjacency with that neighbor.
2. **Exchange Link-State Advertisements** - After adjacencies are established, routers then exchange link-state advertisements (LSAs). LSAs contain the state and cost of each directly connected link. Routers flood their LSAs to adjacent neighbors. Adjacent neighbors receiving the LSA immediately flood the LSA to other directly connected neighbors, until all routers in the area have all LSAs.

3. Build the Topology Table - After LSAs are received, OSPF-enabled routers build the topology table (LSDB) based on the received LSAs. This database eventually holds all the information about the topology of the network.

4. Execute the SPF Algorithm - Routers then execute the SPF algorithm. The SPF algorithm creates the SPF tree.

From the SPF tree, the best paths are inserted into the routing table. Routing decisions are made based on the entries in the routing table.

Single Area and Multiarea OSPF

To make OSPF more efficient and scalable, OSPF supports hierarchical routing using areas. An OSPF area is a group of routers that share the same link-state information in their LSDBs.

OSPF can be implemented in one of two ways:

- **Single-Area OSPF** - All routers are in one area called the backbone area (area 0).
- **Multiarea OSPF** - OSPF is implemented using multiple areas, in a hierarchical fashion. All areas must connect to the backbone area (area 0). Routers interconnecting the areas are referred to as Area Border Routers (ABR).

With multiarea OSPF, OSPF can divide one large autonomous system (AS) into smaller areas, to support hierarchical routing. With hierarchical routing, routing still occurs between the areas (interarea routing), while many of the processor intensive routing operations, such as recalculating the database, are kept within an area.

For instance, any time a router receives new information about a topology change within the area (including the addition, deletion, or modification of a link) the router must rerun the SPF algorithm, create a new SPF tree, and update the routing table. The SPF algorithm is CPU-intensive and the time it takes for calculation depends on the size of the area.

Note: Topology changes are distributed to routers in other areas in a distance vector format. In other words, these routers only update their routing tables and do not need to rerun the SPF algorithm.

Too many routers in one area would make the LSDBs very large and increase the load on the CPU. Therefore, arranging routers into areas effectively partitions a potentially large database into smaller and more manageable databases.

The hierarchical-topology possibilities of multiarea OSPF have these advantages:

- **Smaller routing tables** - Fewer routing table entries because network addresses can be summarized between areas. Route summarization is not enabled by default.
- **Reduced link-state update overhead** - Minimizes processing and memory requirements.
- **Reduced frequency of SPF calculations** - Localizes the impact of a topology change within an area. For instance, it minimizes routing update impact because LSA flooding stops at the area boundary.

Topic 89: OSPF Messages:

OSPF Message Encapsulation

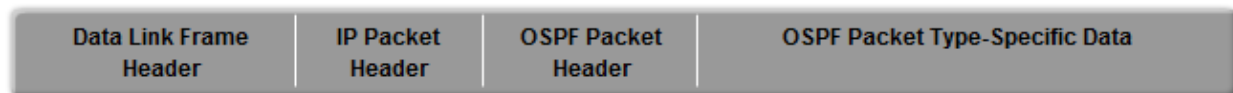
OSPF messages transmitted over an Ethernet link contain the following information:

Data Link Ethernet Frame Header - Identifies the destination multicast MAC addresses 01-00-5E-00-00-05 or 01-00-5E-00-00-06.

IP Packet Header - Identifies the IPv4 protocol field 89 which indicates that this is an OSPF packet. It also identifies one of two OSPF multicast addresses, 224.0.0.5 or 224.0.0.6.

OSPF Packet Header - Identifies the OSPF packet type, the router ID and the area ID.

OSPF Packet Type Specific Data - Contains the OSPF packet type information. The content differs depending on the packet type. In this case, it is an IPv4 Header.



OSPF Hello Packet

Hello Packet

The OSPF Type 1 packet is the Hello packet. Hello packets are used to:

Discover OSPF neighbors and establish neighbor adjacencies.

Advertise parameters on which two routers must agree to become neighbors.

Elect the Designated Router (DR) and Backup Designated Router (BDR) on multiaccess networks like Ethernet and Frame Relay. Point-to-point links do not require DR or BDR.

Fields contained in the Type 1 Hello packet:

Type - Identifies the type of packet. A one (1) indicates a Hello packet. A value 2 identifies a DBD packet, 3 an LSR packet, 4 an LSU packet, and 5 an LSAck packet.

Router ID - A 32-bit value expressed in dotted decimal notation (an IPv4 address) used to uniquely identifying the originating router.

Area ID - Area from which the packet originated.

Network Mask - Subnet mask associated with the sending interface.

Hello Interval - Specifies the frequency, in seconds, at which a router sends Hello packets. The default Hello interval on multiaccess networks is 10 seconds. This timer must be the same on neighboring routers; otherwise, an adjacency is not established.

Router Priority - Used in a DR/BDR election. The default priority for all OSPF routers is 1, but can be manually altered from 0 to 255. The higher the value, the more likely the router becomes the DR on the link.

Dead Interval - Is the time in seconds that a router waits to hear from a neighbor before declaring the neighboring router out of service. By default, the router Dead Interval is four times the Hello interval. This timer must be the same on neighboring routers; otherwise, an adjacency is not established.

Designated Router (DR) - Router ID of the DR.

Backup Designated Router (BDR) - Router ID of the BDR.

List of Neighbors - List that identifies the router IDs of all adjacent routers.

OSPF Hello Interval

OSPF Hello packets are transmitted to multicast address 224.0.0.5 in IPv4 and FF02::5 in IPv6 (all OSPF routers) every:

- 10 seconds (default on multiaccess and point-to-point networks)
- 30 seconds (default on nonbroadcast multiaccess [NBMA] networks; for example, Frame Relay)

The Dead interval is the period that the router waits to receive a Hello packet before declaring the neighbor down. If the Dead interval expires before the routers receive a Hello packet, OSPF removes that neighbor from its LSDB. The router floods the LSDB with information about the down neighbor out all OSPF-enabled interfaces.

Cisco uses a default of 4 times the Hello interval:

- 40 seconds (default on multiaccess and point-to-point networks)
- 120 seconds (default on NBMA networks; for example, Frame Relay)

OSPF Link State Updates

Routers initially exchange Type 2 DBD packets, which is an abbreviated list of the sending router's LSDB and is used by receiving routers to check against the local LSDB.

A Type 3 LSR packet is used by the receiving routers to request more information about an entry in the DBD.

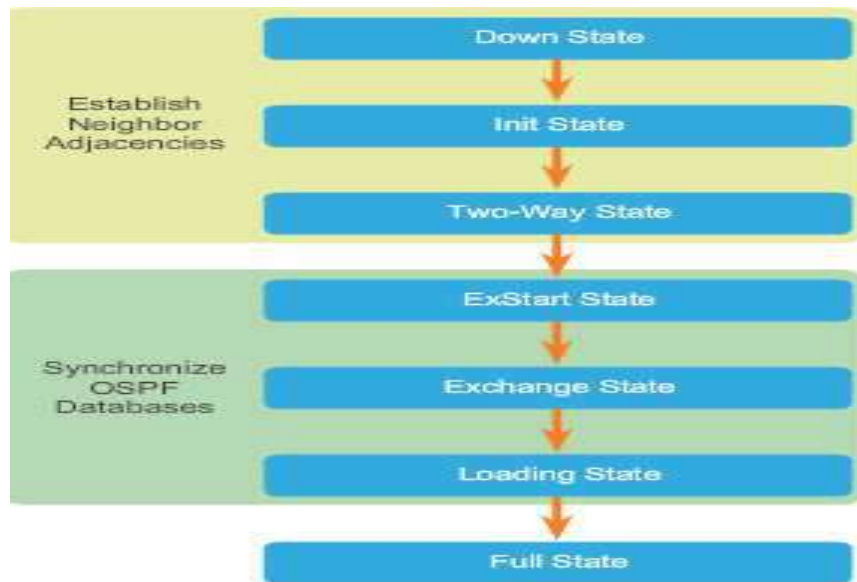
The Type 4 LSU packet is used to reply to an LSR packet.

LSUs are also used to forward OSPF routing updates, such as link changes. Specifically, an LSU packet can contain 11 different types of OSPFv2 LSAs. OSPFv3 renamed several of these LSAs and also contains two additional LSAs.

Note: The difference between the LSU and LSA terms can sometimes be confusing because these terms are often used interchangeably. However, an LSU contains one or more LSAs.

Topic 90: OSPF Operation:

OSPF Operational States



Establish Neighbor Adjacencies

When OSPF is enabled on an interface, the router must determine if there is another OSPF neighbor on the link. To accomplish this, the router forwards a Hello packet that contains its router ID out all OSPF-enabled interfaces. The OSPF router ID is used by the OSPF process to uniquely identify each router in the OSPF area. A router ID is an IP address assigned to identify a specific router among OSPF peers.

When a neighboring OSPF-enabled router receives a Hello packet with a router ID that is not within its neighbor list, the receiving router attempts to establish an adjacency with the initiating router.

The action performed in Two-Way state depends on the type of inter-connection between the adjacent routers:

If the two adjacent neighbors are interconnected over a point-to-point link, then they immediately transition from the Two-Way state to the database synchronization phase.

If the routers are interconnected over a common Ethernet network, then a designated router DR and a BDR must be elected.

Hello packets are continually exchanged to maintain router information.

OSPF DR and BDR

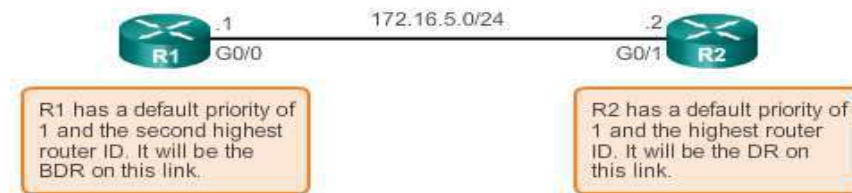
Why is a DR and BDR election necessary?

Multiaccess networks can create two challenges for OSPF regarding the flooding of LSAs:

Creation of multiple adjacencies - Ethernet networks could potentially interconnect many OSPF routers over a common link. Creating adjacencies with every router is unnecessary and undesirable. It would lead to an excessive number of LSAs exchanged between routers on the same network.

Extensive flooding of LSAs - Link-state routers flood their LSAs any time OSPF is initialized, or when there is a change in the topology. This flooding can become excessive.

The solution to managing the number of adjacencies and the flooding of LSAs on a multiaccess network is the DR. On multiaccess networks, OSPF elects a DR to be the collection and distribution point for LSAs sent and received. A BDR is also elected in case the DR fails. All other routers become DROTHERs. A DROTHER is a router that is neither the DR nor the BDR.



Synchronizing OSPF Databases

After the Two-Way state, routers transition to database synchronization states. While the Hello packet was used to establish neighbor adjacencies, the other four types of OSPF packets are used during the process of exchanging and synchronizing LSDBs.

In the ExStart state, a master and slave relationship is created between each router and its adjacent DR and BDR. The router with the higher router ID acts as the master for the Exchange state. In Figure above, R2 becomes the master.

In the Exchange state, the master and slave routers exchange one or more DBD packets. A DBD packet includes information about the LSA entry header that appears in the router's LSDB. The entries can be about a link or about a network. Each LSA entry header includes information about the link-state type, the address of the advertising router, the link's cost, and the sequence number. The router uses the sequence number to determine the newness of the received link-state information.

In Figure below, R2 sends a DBD packet to R1. When R1 receives the DBD, it performs the following actions:

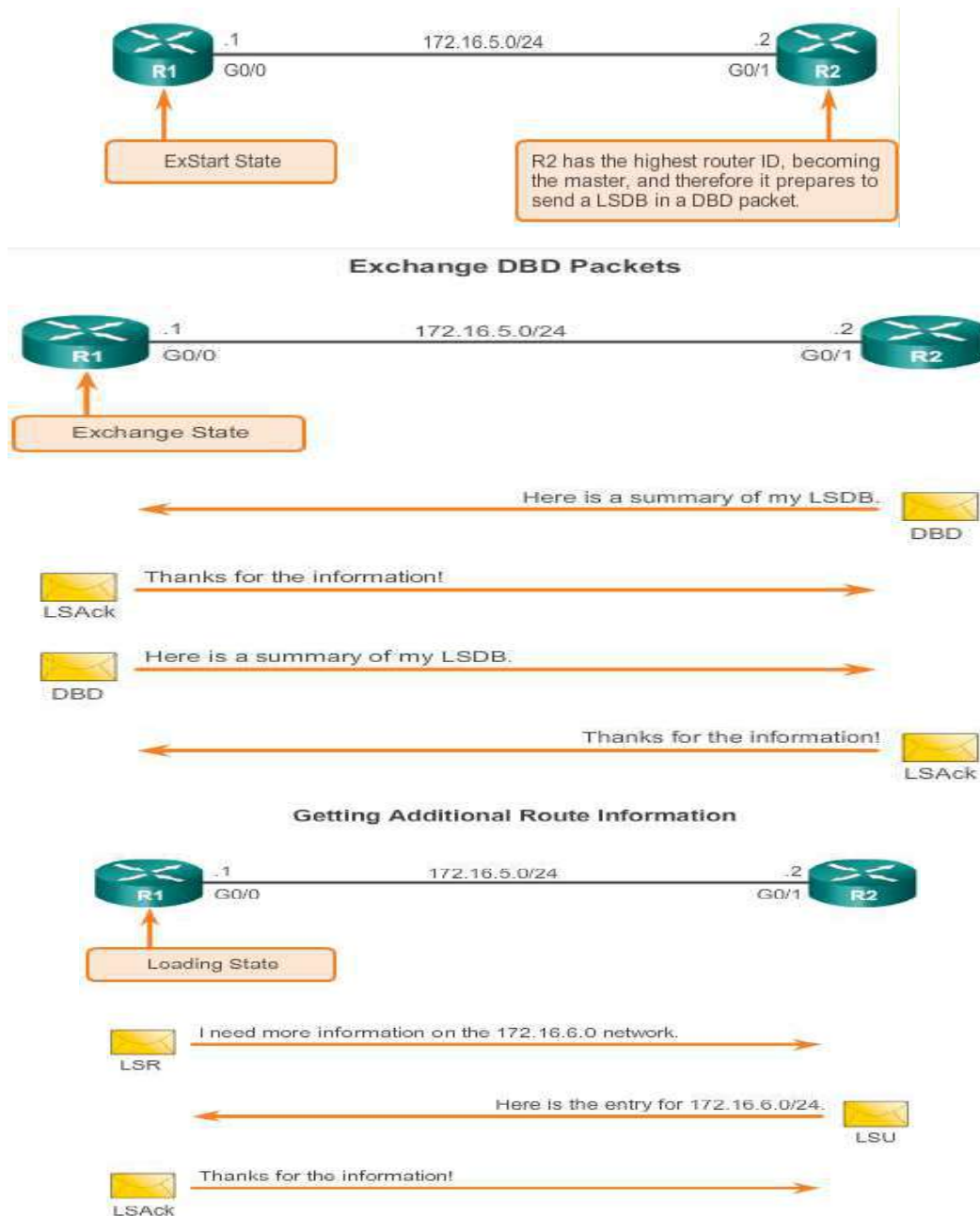
1. It acknowledges the receipt of the DBD using the LSack packet.
2. R1 then sends DBD packets to R2.
3. R2 acknowledges R1.

R1 compares the information received with the information it has in its own LSDB. If the DBD packet has a more current link-state entry, the router transitions to the Loading state.

After all LSRs have been satisfied for a given router, the adjacent routers are considered synchronized and in a full state.

As long as the neighboring routers continue receiving Hello packets, the network in the transmitted LSAs remain in the topology database. After the topological databases are synchronized, updates (LSUs) are sent only to neighbors when:

- A change is perceived (incremental updates)
- Every 30 minutes



Topic 91: Configure Single Area OSPFv2:

Introduced in 1991, OSPFv2 is a link-state routing protocol for IPv4. OSPF was designed as an alternative to another IPv4 routing protocol, RIP.

Every router requires a router ID to participate in an OSPF domain. The router ID can be defined by an administrator or automatically assigned by the router. The router ID is used by the OSPF-enabled router to:

Uniquely identify the router - The router ID is used by other routers to uniquely identify each router within the OSPF domain and all packets that originate from them.

Participate in the election of the DR - In a multiaccess LAN environment, the election of the DR occurs during initial establishment of the OSPF network. When OSPF links become active, the routing device configured with the highest priority is elected the DR. Assuming there is no priority configured, or there is a tie, then the router with the highest router ID is elected the DR. The routing device with the second highest router ID is elected the BDR.

But how does the router determine the router ID? Cisco routers derive the router ID based on one of three criteria, in the following preferential order:

The router ID is explicitly configured using the OSPF **router-id** *rid* router configuration mode command. The *rid* value is any 32-bit value expressed as an IPv4 address. This is the recommended method to assign a router ID.

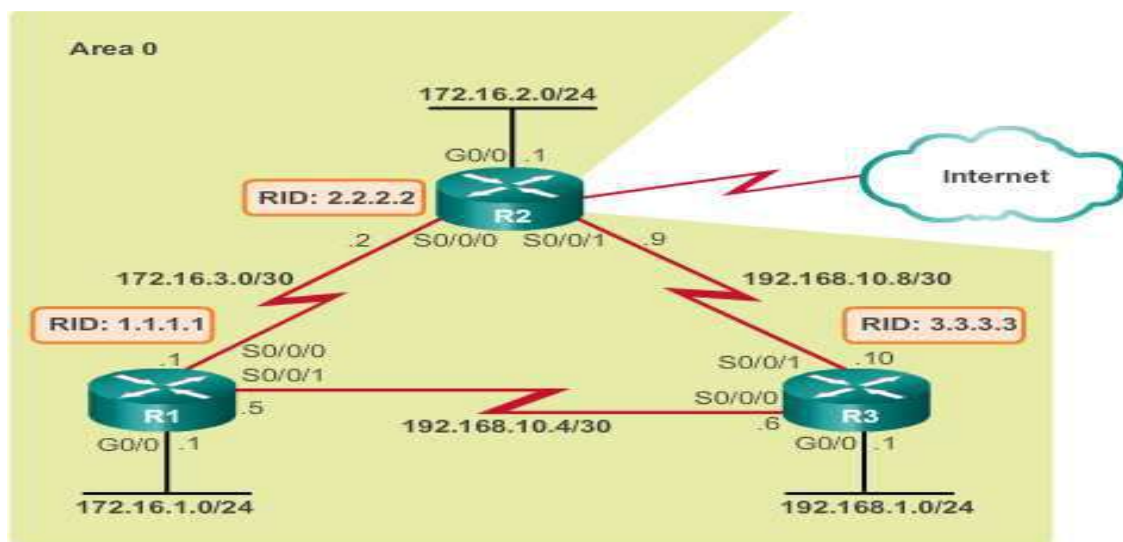
If the router ID is not explicitly configured, the router chooses the highest IPv4 address of any of configured loopback interfaces. This is the next best alternative to assigning a router ID.

If no loopback interfaces are configured, then the router chooses the highest active IPv4 address of any of its physical interfaces. This is the least recommended method because it makes it more difficult for administrators to distinguish between specific routers.

If the router uses the highest IPv4 address for the router ID, the interface does not need to be OSPF-enabled. This means that the interface address does not need to be included in one of the OSPF **network** commands for the router to use that IP address as the router ID. The only requirement is that the interface is active and in the up state.

Note: The router ID looks like an IP address, but it is not routable and, therefore, is not included in the routing table, unless the OSPF routing process chooses an interface (physical or loopback) that is appropriately defined by a **network** command.

Configuring an OSPF Router ID



Use the **router-id** *rid* router configuration mode command to manually assign a 32-bit value expressed as an IPv4 address to a router. An OSPF router identifies itself to other routers using this router ID.

As shown in Figure above, R1 is configured with a router ID of 1.1.1.1, R2 with 2.2.2.2, and R3 with 3.3.3.3.

Use the **show ip protocols** command to verify the router ID.

Sometimes a router ID needs to be changed, for example, when a network administrator establishes a new router ID scheme for the network. However, after a router selects a router ID, an active OSPF router does not allow the router ID to be changed until the router is reloaded or the OSPF process cleared.

Clearing the OSPF process is the preferred method to reset the router ID.

The OSPF routing process is cleared using the **clear ip ospf process** privileged EXEC mode command. This forces OSPF on R1 to transition to the Down and Init states. The **show ip protocols** command verifies that the router ID has changed.

Using a Loopback Interface as the Router ID

A router ID can also be assigned using a loopback interface.

The IPv4 address of the loopback interface should be configured using a 32-bit subnet mask (255.255.255.255). This effectively creates a host route. A 32-bit host route does not get advertised as a route to other OSPF routers.

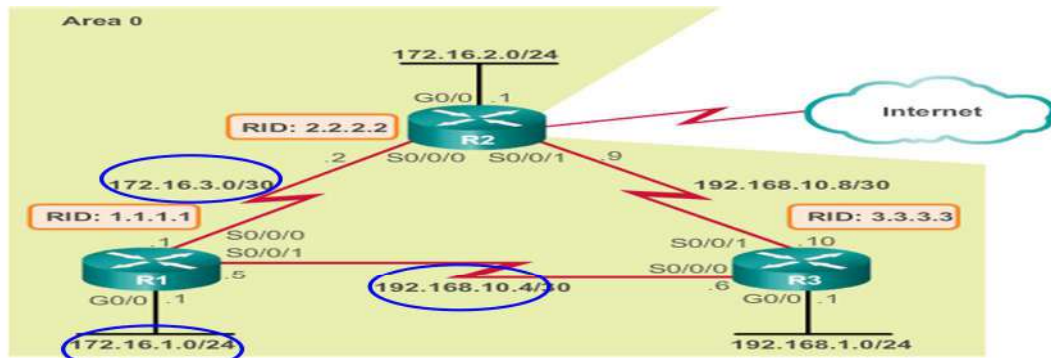
Enabling OSPF on Interfaces

The **network** command determines which interfaces participate in the routing process for an OSPF area. Any interfaces on a router that match the network address in the **network** command are enabled to send and receive OSPF packets. As a result, the network (or subnet) address for the interface is included in OSPF routing updates.

The basic command syntax is **network** *network-address wildcard-mask* **area** *area-id*.

The **area** *area-id* syntax refers to the OSPF area. When configuring single-area OSPF, the **network** command must be configured with the same *area-id* value on all routers. Although any area ID can be used, it is good practice to use an area ID of 0 with single-area OSPF. This convention makes it easier if the network is later altered to support multiarea OSPF.

The figure below displays the reference topology.



```

R1(config)# router ospf 10
R1(config-router)# route-id 1.1.1.1
R1(config-router)# network 172.16.1.0 0.0.0.255 area 0
R1(config-router)# network 172.16.3.0 0.0.0.3 area 0
R1(config-router)# network 192.168.10.4 0.0.0.3 area 0
R1(config-router)# end

```

As an alternative, OSPFv2 can be enabled using the **network** *intf-ip-address* **0.0.0.0** *area area-id* router configuration mode command.

The advantage of specifying the interface is that the wildcard mask calculation is not necessary. OSPFv2 uses the interface address and subnet mask to determine the network to advertise.

Topic 92: Packet Tracer – Configure Basic Single Area OSPFv2 - 1:

- Build Network
- Configure Devices
- Configure OSPF Routing
- Verify OSPF Routing
- Change Router ID Assignments
- Configure OSPF Passive Interfaces
- Change OSPF Metrics

Topic 93: Packet Tracer – Configure Basic Single Area OSPFv2 - 2:

This is the continuation of the previous topic.

Topic 94: Packet Tracer – Configure Basic Single Area OSPFv2 - 3:

This is the continuation of the previous topic.

Topic 95: Packet Tracer – Configure Basic Single Area OSPFv2 - 4:

This is the continuation of the previous topic.

Topic 96: Packet Tracer – Configuring OSPFv2 in a Single Area:

- Configure OSPFv2 Routing
- Verify Configurations

Topic 97: OSPFv2 Cost:

Recall that a routing protocol uses a metric to determine the best path of a packet across a network. A metric gives indication of the overhead that is required to send packets across a certain interface. OSPF uses cost as a metric. A lower cost indicates a better path than a higher cost.

The cost of an interface is inversely proportional to the bandwidth of the interface. Therefore, a higher bandwidth indicates a lower cost. More overhead and time delays equal a higher cost. Therefore, a 10-Mb/s Ethernet line has a higher cost than a 100-Mb/s Ethernet line.

The formula used to calculate the OSPF cost is:

$$\text{Cost} = \text{reference bandwidth} / \text{interface bandwidth}$$

The default reference bandwidth is 10^8 (100,000,000); therefore, the formula is:

$$\text{Cost} = 100,000,000 \text{ bps} / \text{interface bandwidth in bps}$$

Refer to the table in the figure below for a breakdown of the cost calculation. Notice that FastEthernet, Gigabit Ethernet, and 10 GigE interfaces share the same cost, because the OSPF cost value must be an integer. Consequently, because the default reference bandwidth is set to 100 Mb/s, all links that are faster than Fast Ethernet also have a cost of 1.

Interface Type	Reference Bandwidth in bps	Default Bandwidth in bps	Cost
10 Gigabit Ethernet 10 Gbps	1,000,000,000	÷ 10,000,000,000	1
Gigabit Ethernet 1 Gbps	1,000,000,000	÷ 1,000,000,000	1
Fast Ethernet 100 Mbps	1,000,000,000	÷ 100,000,000	10
Ethernet 10 Mbps	1,000,000,000	÷ 10,000,000	100
Serial 1.544 Mbps	1,000,000,000	÷ 1,544,000	647
Serial 128 kbps	1,000,000,000	÷ 128,000	7812
Serial 64 kbps	1,000,000,000	÷ 64,000	15625

The cost of an OSPF route is the accumulated value from one router to the destination network.

All interfaces have default bandwidth values assigned to them. As with reference bandwidth, interface bandwidth values do not actually affect the speed or capacity of the link. Instead, they are used by OSPF to compute the routing metric. Therefore, it is important that the bandwidth value reflect the actual speed of the link so that the routing table has accurate best path information.

Although the bandwidth values of Ethernet interfaces usually match the link speed, some other interfaces may not. For instance, the actual speed of serial interfaces is often different than the default bandwidth. On Cisco routers, the default bandwidth on most serial interfaces is set to 1.544 Mb/s.

Use the show interfaces command to view the interface bandwidth setting.

As an alternative to setting the default interface bandwidth, the cost can be manually configured on an interface using the **ip ospf cost value** interface configuration command.

An advantage of configuring a cost over setting the interface bandwidth is that the router does not have to calculate the metric when the cost is manually configured. In contrast, when the interface bandwidth is configured, the router must calculate the OSPF cost based on the bandwidth. The **ip ospf cost** command is useful in multi-vendor environments where non-Cisco routers may use a metric other than bandwidth to calculate the OSPF costs.

Both the **bandwidth** interface command and the **ip ospf cost** interface command achieve the same result, which is to provide an accurate value for use by OSPF in determining the best route.

Topic 98: Verify OSPFv2:

Figure below shows the reference topology.

Use the **show ip ospf neighbor** command to verify that the router has formed an adjacency with its neighboring routers. If the router ID of the neighboring router is not displayed, or if it does not show as being in a state of FULL, the two routers have not formed an OSPF adjacency.

If two routers do not establish adjacency, link-state information is not exchanged. Incomplete LSDBs can cause inaccurate SPF trees and routing tables. Routes to destination networks may not exist, or may not be the most optimum path.

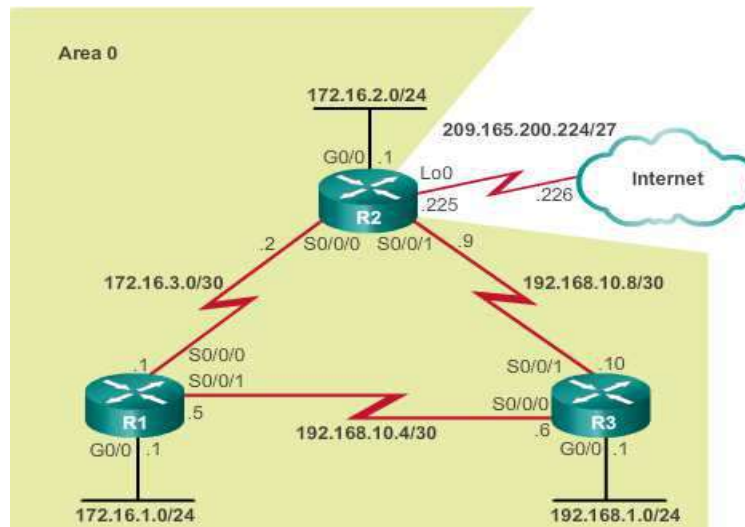


Figure below displays the neighbor adjacency of R1. For each neighbor, this command displays the following output:

Neighbor ID - The router ID of the neighboring router.

Pri - The OSPF priority of the interface. This value is used in the DR and BDR election.

State - The OSPF state of the interface. FULL state means that the router and its neighbor have identical OSPF LSDBs. On multiaccess networks, such as Ethernet, two routers that are adjacent may have their

states displayed as 2WAY. The dash indicates that no DR or BDR is required because of the network type.

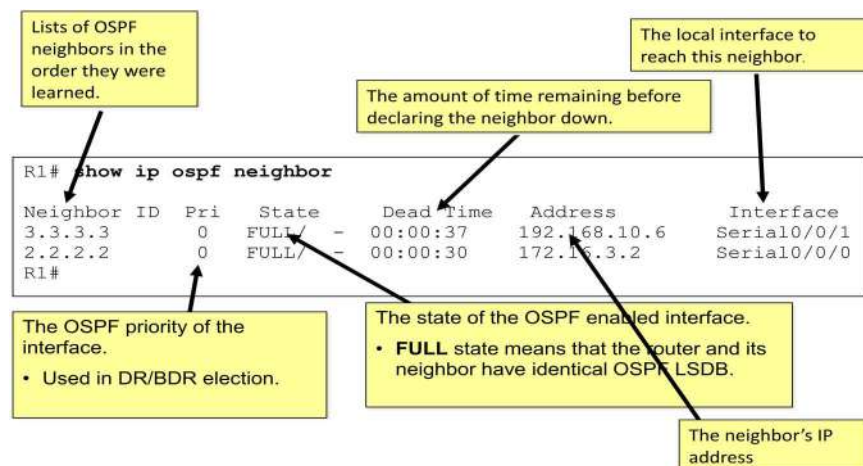
Dead Time - The amount of time remaining that the router waits to receive an OSPF Hello packet from the neighbor before declaring the neighbor down. This value is reset when the interface receives a Hello packet.

Address - The IPv4 address of the neighbor's interface to which this router is directly connected.

Interface - The interface on which this router has formed adjacency with the neighbor.

Two routers may not form an OSPF adjacency if:

- The subnet masks do not match, causing the routers to be on separate networks.
- OSPF Hello or Dead Timers do not match.
- OSPF Network Types do not match.
- There is a missing or incorrect OSPF **network** command.



Verify OSPF Interface Settings

The quickest way to verify OSPF interface settings is to use the **show ip ospf interface** command. This command provides a detailed list for every OSPF-enabled interface. The command is useful to determine whether the **network** statements were correctly composed.

To get a summary of OSPF-enabled interfaces, use the **show ip ospf interface brief**.

```
R1# show ip ospf interface brief
Interface PID Area IP Address/Mask Cost State Nbrs F/C
Se0/0/1 10 0 192.168.10.5/30 15625 P2P 1/1
Se0/0/0 10 0 172.16.3.1/30 647 P2P 1/1
Gi0/0 10 0 172.16.1.1/24 1 DR 0/0
R1#
```

Topic 99: OSPFv3 vs. OSPFv2:

- OSPFv3 is the OSPFv2 equivalent for exchanging IPv6 prefixes. Recall that in IPv6, the network address is referred to as the prefix and the subnet mask is called the prefix-length.
- Similar to its IPv4 counterpart, OSPFv3 exchanges routing information to populate the IPv6 routing table with remote prefixes.
- OSPFv2 runs over the IPv4 network layer, communicating with other OSPF IPv4 peers, and advertising only IPv4 routes.
- OSPFv3 has the same functionality as OSPFv2, but uses IPv6 as the network layer transport, communicating with OSPFv3 peers and advertising IPv6 routes. OSPFv3 also uses the SPF algorithm as the computation engine to determine the best paths throughout the routing domain.
- As with all IPv6 routing protocols, OSPFv3 has separate processes from its IPv4 counterpart. The processes and operations are basically the same as in the IPv4 routing protocol, but run independently. OSPFv2 and OSPFv3 each have separate adjacency tables, OSPF topology tables, and IP routing tables.
- The OSPFv3 configuration and verification commands are similar those used in OSPFv2.

OSPFv2 and OSPFv3 Similarities

- **Link-state** - OSPFv2 and OSPFv3 are both classless link-state routing protocols.
- **Routing algorithm** - OSPFv2 and OSPFv3 use the SPF algorithm to make routing decisions.
- **Metric** - The RFCs for both OSPFv2 and OSPFv3 define the metric as the cost of sending packets out the interface. OSPFv2 and OSPFv3 can be modified using the **auto-cost reference-bandwidth ref-bw** router configuration mode command. The command only influences the OSPF metric where it was configured. For example, if this command was entered for OSPFv3, it does not affect the OSPFv2 routing metrics.
- **Areas** - The concept of multiple areas in OSPFv3 is the same as in OSPFv2. Multiareas that minimize link-state flooding and provide better stability with the OSPF domain.
- **OSPF packet types** - OSPFv3 uses the same five basic packet types as OSPFv2 (Hello, DBD, LSR, LSU, and LSAck).
- **Neighbor discovery mechanism** - The neighbor state machine, including the list of OSPF neighbor states and events, remains unchanged. OSPFv2 and OSPFv3 use the Hello mechanism to learn about neighboring routers and form adjacencies. However, in OSPFv3, there is no requirement for matching subnets to form neighbor adjacencies. This is because neighbor adjacencies are formed using link-local addresses, not global unicast addresses.
- **DR/BDR election process** - The DR/BDR election process remains unchanged in OSPFv3.
- **Router ID** - Both OSPFv2 and OSPFv3 use a 32-bit number for the router ID represented in dotted-decimal notation. Typically this is an IPv4 address. The OSPF **router-id** command must be used to configure the router ID. The process in determining the 32-bit Router ID is the same in both protocols. Use an explicitly-configured router ID; otherwise, the highest loopback IPv4 address becomes the router ID.

OSPFv2 vs. OSPFv3

The differences between OSPFv2 and OSPFv3:

- **Advertises** - OSPFv2 advertises IPv4 routes, whereas OSPFv3 advertises routes for IPv6.
- **Source address** - OSPFv2 messages are sourced from the IPv4 address of the exit interface. In OSPFv3, OSPF messages are sourced using the link-local address of the exit interface.
- **All OSPF router multicast addresses**- OSPFv2 uses 224.0.0.5; whereas, OSPFv3 uses FF02::5.
- **DR/BDR multicast address** - OSPFv2 uses 224.0.0.6; whereas, OSPFv3 uses FF02::6.
- **Advertise networks** - OSPFv2 advertises networks using the **network** router configuration command; whereas, OSPFv3 uses the **ipv6 ospf process-id area area-id** interface configuration command.
- **IP unicast routing** - Enabled, by default, in IPv4; whereas, the **ipv6 unicast-routing** global configuration command must be configured.
- **Authentication** - OSPFv2 uses either plaintext authentication or MD5 authentication. OSPFv3 uses IPv6 authentication.

Topic 100: Configure OSPFv3:

OSPFv3 Network Topology

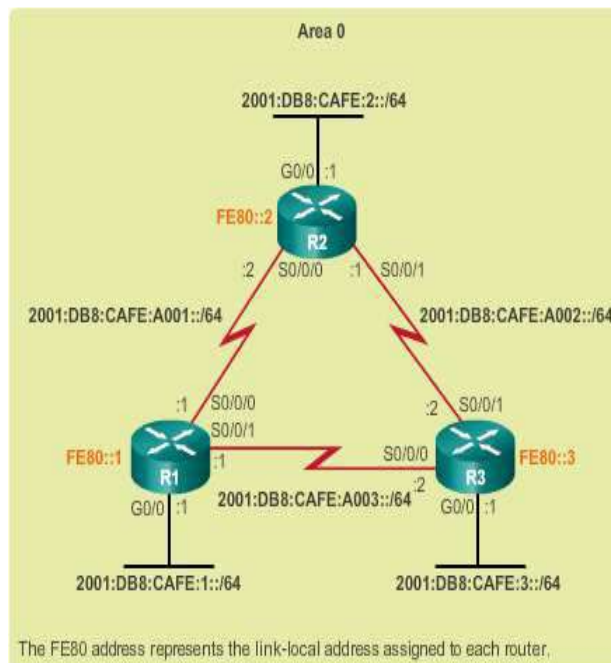


Figure above displays the network topology that is used to configure OSPFv3.

In this topology, none of the routers have IPv4 addresses configured. A network with router interfaces configured with IPv4 and IPv6 addresses is referred to as dual-stacked. A dual-stacked network can have OSPFv2 and OSPFv3 simultaneously-enabled.

The steps to configure basic OSPFv3 in a single area.

Link Local Addresses

The output of the **show ipv6 interface brief** command confirms that the correct global IPv6 addresses have been successfully configured and that the interfaces are enabled.

Link-local addresses are automatically created when an IPv6 global unicast address is assigned to the interface. Global unicast addresses are not required on an interface; however, IPv6 link-local addresses are.

Unless configured manually, Cisco routers create the link-local address using FE80::/10 prefix and the EUI-64 process. EUI-64 involves using the 48-bit Ethernet MAC address, inserting FFFE in the middle and flipping the seventh bit. For serial interfaces, Cisco uses the MAC address of an Ethernet interface.

Configuring Link Local Addresses on R1

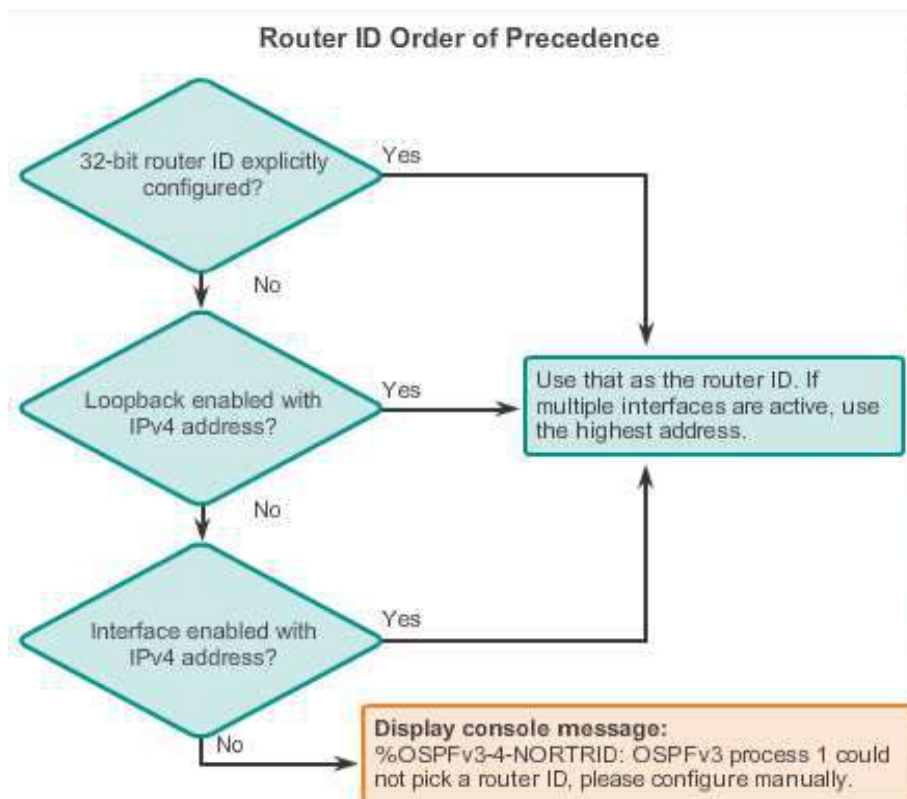
Link-local addresses created using the EUI-64 format or in some cases, random interface IDs, make it difficult to recognize and remember those addresses. Because IPv6 routing protocols use IPv6 link-local addresses for unicast addressing and next-hop address information in the routing table, it is common practice to make it an easily recognizable address.

Configuring the link-local address manually provides the ability to create an address that is recognizable and easier to remember. As well, a router with several interfaces can assign the same link-local address to each IPv6 interface. This is because the link-local address is only required for local communications.

Link-local addresses can be configured manually using the same interface command used to create IPv6 global unicast addresses, but appending the **link-local** keyword to the **ipv6 address** command.

A link-local address has a prefix within the range FE80 to FEBF. When an address begins with this hextet (16-bit segment) the **link-local** keyword must follow the address.

Configuring OSPFv3 Router ID



Topic 101: Verify OSPFv3:

Verify OSPFv3 Neighbors & Protocol Settings

```
R1# show ipv6 ospf neighbor

OSPFv3 Router with ID (1.1.1.1) (Process ID 10)

Neighbor ID    Pri   State   Dead Time   Interface ID  Interface
3.3.3.3        0    FULL/  -   00:00:39    6             Serial0/0/1
2.2.2.2        0    FULL/  -   00:00:36    6             Serial0/0/0
R1#
```

```
R1# show ipv6 protocols
IPv6 Routing Protocol is "connected"
IPv6 Routing Protocol is "ND"
IPv6 Routing Protocol is "ospf 10"
  Router ID 1.1.1.1
  Number of areas: 1 normal, 0 stub, 0 nssa
  Interfaces (Area 0):
    Serial0/0/1
    Serial0/0/0
    GigabitEthernet0/0
  Redistribution:
    None
R1#
```

Use the **show ipv6 ospf neighbor** command to verify that the router has formed an adjacency with its neighboring routers. If the router ID of the neighboring router is not displayed, or if it does not show as being in a state of FULL, the two routers have not formed an OSPF adjacency.

If two routers do not establish a neighbor adjacency, link-state information is not exchanged. Incomplete LSDBs can cause inaccurate SPF trees and routing tables. Routes to destination networks may not exist or may not be the most optimum path.

Figure above displays the neighbor adjacency of R1. For each neighbor, this command displays the following output:

- **Neighbor ID** - The router ID of the neighboring router.
- **Pri** - The OSPF priority of the interface. Value is used in the DR and BDR election.
- **State** - The OSPF state of the interface. FULL state means that the router and its neighbor have identical OSPF LSDBs. On multiaccess networks such as Ethernet, two routers that are adjacent may have their states displayed as 2WAY. The dash indicates that no DR or BDR is required because of the network type.
- **Dead Time** - The amount of time remaining that the router waits to receive an OSPF Hello packet from the neighbor before declaring the neighbor down. This value is reset when the interface receives a Hello packet.
- **Interface ID** - The interface ID or link ID.
- **Interface** - The interface on which this router has formed adjacency with the neighbor.

As shown in Figure above, the **show ipv6 protocols** command is a quick way to verify vital OSPFv3 configuration information, including the OSPF process ID, the router ID, and the interfaces enabled for OSPFv3.

Topic 102: Packet Tracer – Configuring Basic OSPFv3:

- Configure OSPFv3 Routing
- Verify Connectivity

Topic 103: Multiarea OSPF:

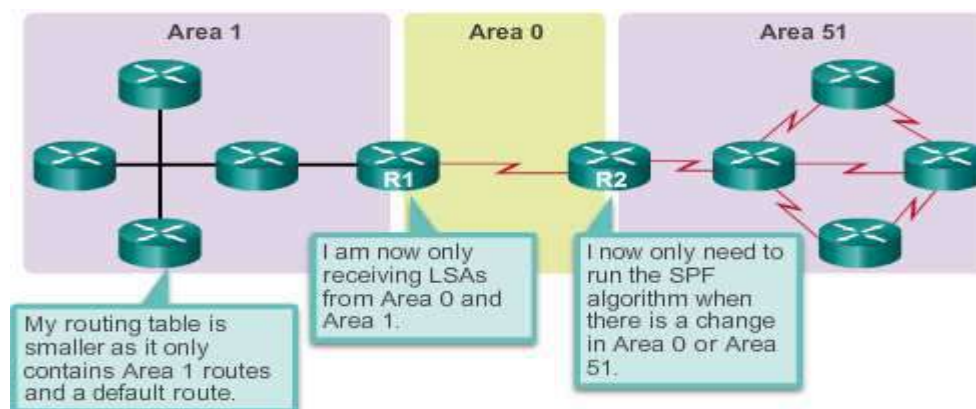
When a large OSPF area is divided into smaller areas, this is called multiarea OSPF. Multiarea OSPF is useful in larger network deployments to reduce processing and memory overhead.

For instance, any time a router receives new information about the topology, as with additions, deletions, or modifications of a link, the router must rerun the SPF algorithm, create a new SPF tree, and update the routing table. The SPF algorithm is CPU-intensive and the time it takes for calculation depends on the size of the area. Too many routers in one area make the LSDB larger and increase the load on the CPU. Therefore, arranging routers into areas effectively partitions one potentially large database into smaller and more manageable databases.

Multiarea OSPF requires a hierarchical network design. The main area is called the backbone area (area 0) and all other areas must connect to the backbone area. With hierarchical routing, routing still occurs between the areas (interarea routing); while many of the tedious routing operations, such as recalculating the database, are kept within an area.

The hierarchical-topology possibilities of multiarea OSPF have these advantages:

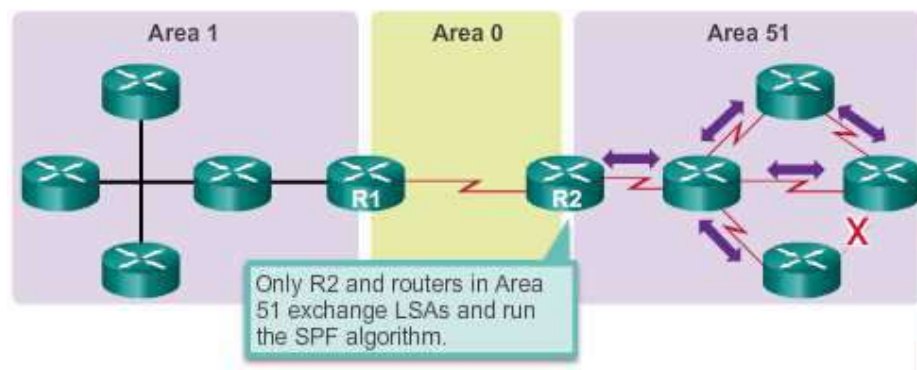
Smaller routing tables - There are fewer routing table entries as network addresses can be summarized between areas. For example in figure below, R1 summarizes the routes from area 1 to area 0 and R2 summarizes the routes from area 51 to area 0. R1 and R2 also propagate a default static route to area 1 and area 51.



Reduced link-state update overhead- Minimizes processing and memory requirements, because there are fewer routers exchanging LSAs.

Reduced frequency of SPF calculations - Localizes impact of a topology change within an area. For instance, it minimizes routing update impact, because LSA flooding stops at the area boundary.

In Figure below, assume a link fails between two internal routers in area 51. Only the routers in area 51 exchange LSAs and rerun the SPF algorithm for this event. R1 does not receive LSAs from area 51 and does not recalculate the SPF algorithm.



OSPF Two-Layer Area Hierarchy

Multiarea OSPF is implemented in a two-layer area hierarchy:

Backbone (Transit) area - An OSPF area whose primary function is the fast and efficient movement of IP packets. Backbone areas interconnect with other OSPF area types. Generally, end users are not found within a backbone area. The backbone area is also called OSPF area 0. Hierarchical networking defines area 0 as the core to which all other areas directly connect.

Regular (Non-backbone) area -Connects users and resources. Regular areas are usually set up along functional or geographical groupings. By default, a regular area does not allow traffic from another area to use its links to reach other areas. All traffic from other areas must cross a transit area.

OSPF enforces this rigid two-layer area hierarchy. The underlying physical connectivity of the network must map to the two-layer area structure, with all non-backbone areas attaching directly to area 0. All traffic moving from one area to another area must traverse the backbone area. This traffic is referred to as interarea traffic.

The optimal number of routers per area varies based on factors such as network stability, but Cisco recommends the following guidelines:

- An area should have no more than 50 routers.
- A router should not be in more than three areas.
- Any single router should not have more than 60 neighbors.

Types of OSPF Routers

OSPF routers of different types control the traffic that goes in and out of areas. The OSPF routers are categorized based on the function they perform in the routing domain.

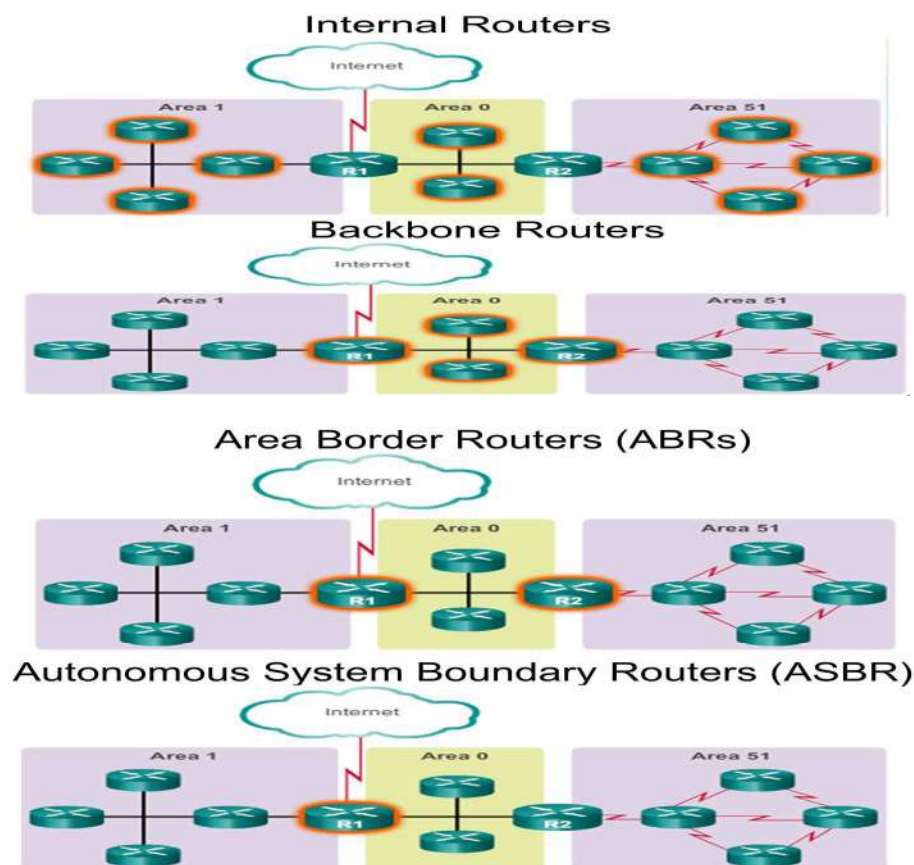
There are four different types of OSPF routers:

- **Internal router** This is a router that has all of its interfaces in the same area. All internal routers in an area have identical LSDBs.

- **Backbone router** This is a router in the backbone area. Generally, the backbone area is set to area 0.
- **Area Border Router (ABR)** This is a router that has interfaces attached to multiple areas. It must maintain separate LSDBs for each area it is connected to, and can route between areas. ABRs are exit points for the area, which means that routing information destined for another area can get there only via the ABR of the local area. ABRs can be configured to summarize the routing information from the LSDBs of their attached areas. ABRs distribute the routing information into the backbone. The backbone routers then forward the information to the other ABRs. In a multiarea network, an area can have one or more ABRs.
- **Autonomous System Boundary Router (ASBR)** This is a router that has at least one interface attached to an external internetwork (another autonomous system), such as a non-OSPF network. An ASBR can import non-OSPF network information to the OSPF network, and vice versa, using a process called route redistribution.

Redistribution in multiarea OSPF occurs when an ASBR connects different routing domains (e.g., EIGRP and OSPF) and configures them to exchange and advertise routing information between those routing domains.

A router can be classified as more than one router type. For example, if a router connects to area 0 and area 1, and in addition, maintains routing information for another, non-OSPF network, it falls under three different classifications: a backbone router, an ABR, and an ASBR.



Topic 104: Multiarea OSPF LSA Operation:

OSPF LSA Type

LSAs are the building blocks of the OSPF LSDB. Individually, they act as database records and provide specific OSPF network details. In combination, they describe the entire topology of an OSPF network or area.

The RFCs for OSPF currently specify up to 11 different LSA types. However, any implementation of multiarea OSPF must support the first five LSAs: LSA 1 to LSA 5. The focus of this topic is on these first five LSAs.

Each router link is defined as an LSA type. The LSA includes a link ID field that identifies, by network number and mask, the object to which the link connects. Depending on the type, the link ID has different meanings. LSAs differ on how they are generated and propagated within the routing domain.

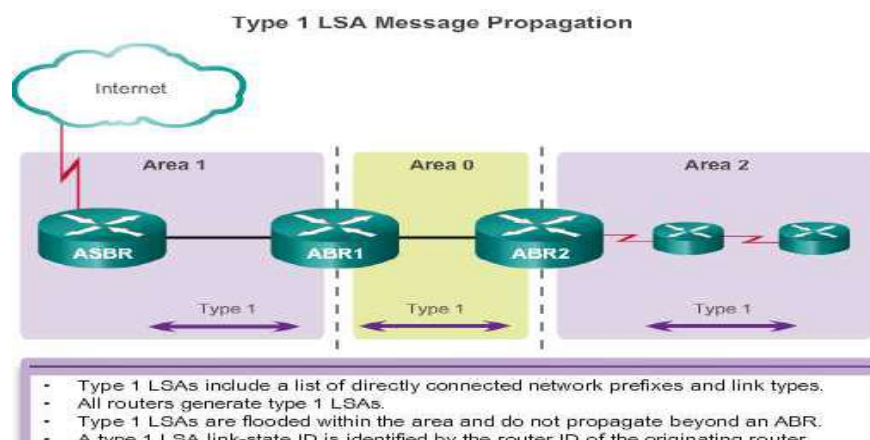
Note: OSPFv3 includes additional LSA types.

LSA Type	Description
1	Router LSA
2	Network LSA
3 and 4	Summary LSAs
5	AS External LSA
6	Multicast OSPF LSA
7	Defined for NSSAs
8	External Attributes LSA for Border Gateway Protocol (BGP)
9, 10, or 11	Opaque LSAs

Most Common LSA Types

LSA Type	Description
1	Router LSA
2	Network LSA
3 and 4	Summary LSAs
5	AS External LSA

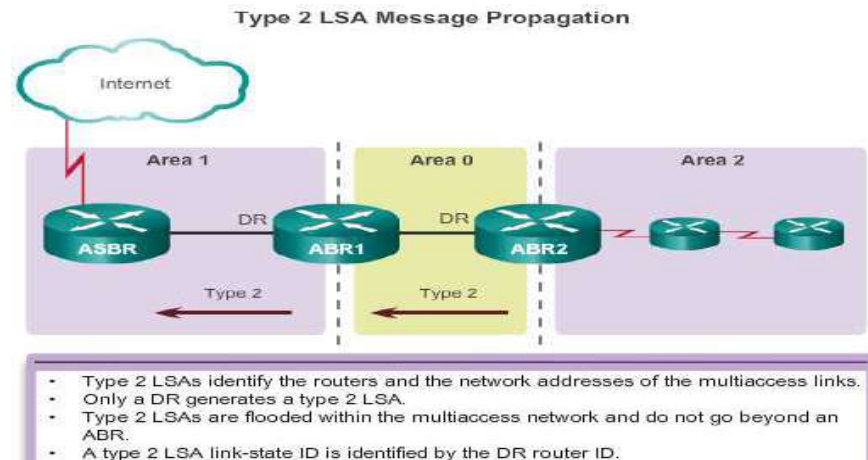
OSPF LSA Type 1



As shown in the figure above, all routers advertise their directly connected OSPF-enabled links in a type 1 LSA and forward their network information to OSPF neighbors. The LSA contains a list of the directly connected interfaces, link types, and link states.

- Type 1 LSAs are also referred to the router link entries.
- Type 1 LSAs are flooded only within the area in which they originated. ABRs subsequently advertise the networks learned from the type 1 LSAs to other areas as type 3 LSAs.
- The type 1 LSA link ID is identified by the router ID of the originating router.

OSPF LSA Type 2



A type 2 LSA only exists for multiaccess and non-broadcast multiaccess (NBMA) networks where there is a DR elected and at least two routers on the multiaccess segment. The type 2 LSA contains the router ID and IP address of the DR, along with the router ID of all other routers on the multiaccess segment. A type 2 LSA is created for every multiaccess network in the area.

The purpose of an a type 2 LSA is to give other routers information about multiaccess networks within the same area.

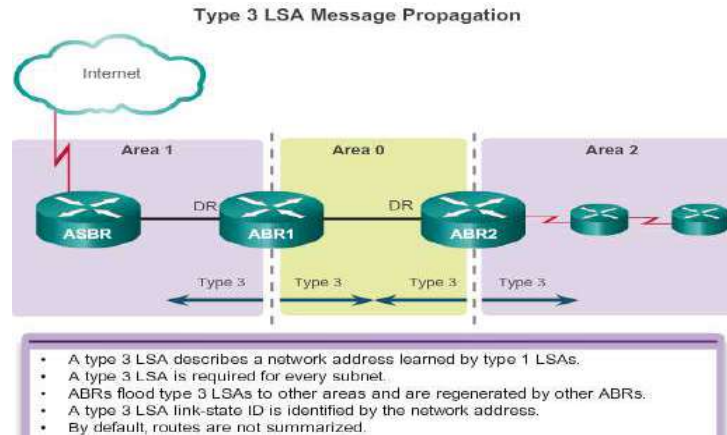
The DR floods type 2 LSAs only within the area in which they originated. Type 2 LSAs are not forwarded outside of an area.

Type 2 LSAs are also referred to the network link entries.

As shown in the figure above, ABR1 is the DR for the Ethernet network in area 1. It generates the type 2 LSA and forwards it into area 1. ABR2 is the DR for the multiaccess network in area 0. There are no multiaccess networks in area 2 and therefore, no type 2 LSAs are ever propagated in that area.

The link-state ID for a network LSA is the IP interface address of the DR that advertises it.

OSPF LSA Type 3



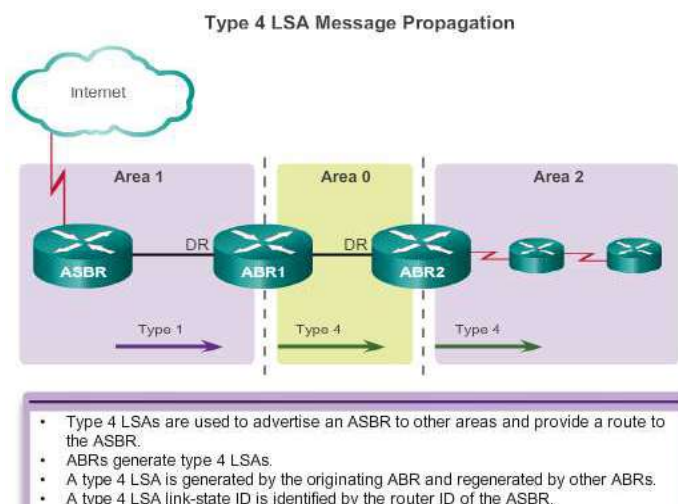
Type 3 LSAs are used by ABRs to advertise networks from other areas. ABRs collect type 1 LSAs in the LSDB. After an OSPF area has converged, the ABR creates a type 3 LSA for each of its learned OSPF networks. Therefore, an ABR with many OSPF routes must create type 3 LSAs for each network.

As shown in the figure above, ABR1 and ABR2 flood type 3 LSAs from one area to other areas. The ABRs propagate the type 3 LSAs into other areas. In a large OSPF deployment with many networks, propagating type 3 LSAs can cause significant flooding problems. For this reason, it is strongly recommended that manual route summarization be configured on the ABR.

The link-state ID is set to the network number and the mask is also advertised.

Receiving a type 3 LSA into its area does not cause a router to run the SPF algorithm. The routes being advertised in the type 3 LSAs are appropriately added to or deleted from the router's routing table, but a full SPF calculation is not necessary.

OSPF LSA Type 4



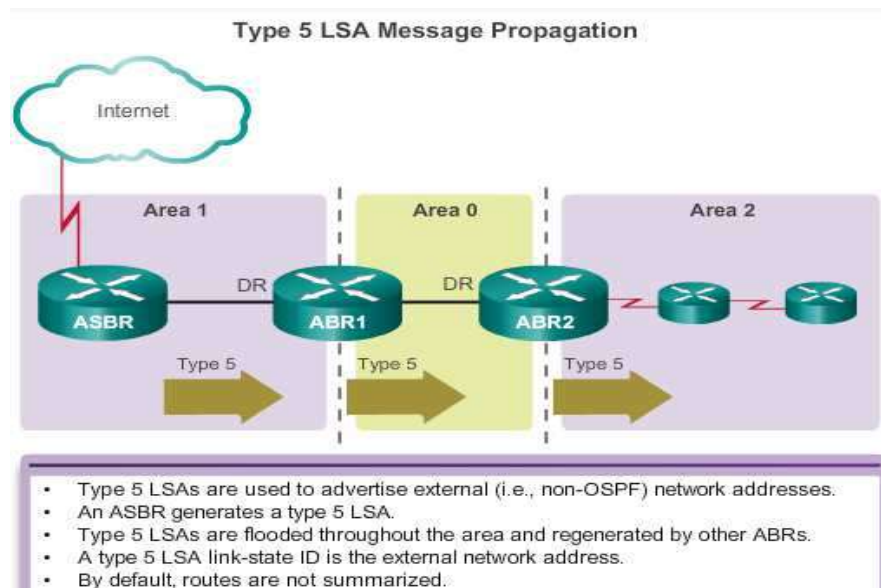
Type 4 and type 5 LSAs are used collectively to identify an ASBR and advertise external networks into an OSPF routing domain.

A type 4 summary LSA is generated by an ABR only when an ASBR exists within an area. A type 4 LSA identifies the ASBR and provides a route to it. All traffic destined to an external autonomous system requires routing table knowledge of the ASBR that originated the external routes.

As shown in the figure above, the ASBR sends a type 1 LSA, identifying itself as an ASBR. The LSA includes a special bit known as the external bit (e bit) that is used to identify the router as an ASBR. When ABR1 receives the type 1 LSA, it notices the e bit, it builds a type 4 LSA, and then floods the type 4 LSA to the backbone (area 0). Subsequent ABRs flood the type 4 LSA into other areas.

The link-state ID is set to the ASBR router ID.

OSPF LSA Type 5



Type 5 external LSAs describe routes to networks outside the OSPF autonomous system. Type 5 LSAs are originated by the ASBR and are flooded to the entire autonomous system.

Type 5 LSAs are also referred to as autonomous system external LSA entries.

In the figure above, the ASBR generates type 5 LSAs for each of its external routes and floods it into the area. Subsequent ABRs also flood the type 5 LSA into other areas. Routers in other areas use the information from the type 4 LSA to reach the external routes.

In a large OSPF deployment with many networks, propagating multiple type 5 LSAs can cause significant flooding problems. For this reason, it is strongly recommended that manual route summarization be configured on the ASBR.

The link-state ID is the external network number.

Topic 105: OSPF Routing Table and Types of Routes:

OSPF Routing Table Entries

```

R1# show ip route
Codes: L - local, C - connected, S - static, R - RIP, M - mobile, B - BGP
       D - EIGRP, EX - EIGRP external, O - OSPF, IA - OSPF inter area
       N1 - OSPF NSSA external type 1, N2 - OSPF NSSA external type 2
       E1 - OSPF external type 1, E2 - OSPF external type 2
       i - IS-IS, su - IS-IS summary, L1 - IS-IS level-1, L2 - IS-IS level-2
       ia - IS-IS inter area, * - candidate default, U - per-user static route
       o - ODR, P - periodic downloaded static route, R - NHRP, l - LISP
       + - replicated route, % - next hop override

Gateway of last resort is 192.168.10.2 to network 0.0.0.0

O*E2 0.0.0.0/0 [110/1] via 192.168.10.2, 00:00:19, Serial0/0/0
    10.0.0.0/8 is variably subnetted, 5 subnets, 2 masks
    C 10.1.1.0/24 is directly connected, GigabitEthernet0/0
    L 10.1.1.1/32 is directly connected, GigabitEthernet0/0
    C 10.1.2.0/24 is directly connected, GigabitEthernet0/1
    L 10.1.2.1/32 is directly connected, GigabitEthernet0/1
    O 10.2.1.0/24 [110/648] via 192.168.10.2, 00:04:34, Serial0/0/0
    O IA 192.168.1.0/24 [110/1295] via 192.168.10.2, 00:01:48, Serial0/0/0
    O IA 192.168.2.0/24 [110/1295] via 192.168.10.2, 00:01:48, Serial0/0/0
    192.168.10.0/24 is variably subnetted, 3 subnets, 2 masks
    C 192.168.10.0/30 is directly connected, Serial0/0/0
    L 192.168.10.1/32 is directly connected, Serial0/0/0
    O 192.168.10.4/30 [110/1294] via 192.168.10.2, 00:01:55, Serial0/0/0
R1#

```

Figure above provides a sample routing table for a multiarea OSPF topology with a link to an external non-OSPF network. OSPF routes in an IPv4 routing table are identified using the following descriptors:

- - Router (type 1) and network (type 2) LSAs describe the details within an area. The routing table reflects this link-state information with a designation of **O**, meaning that the route is intra-area.
- **IA** - When an ABR receives summary LSAs, it adds them to its LSDB and regenerates them into the local area. When an ABR receives external LSAs, it adds them to its LSDB and floods them into the area. The internal routers then assimilate the information into their databases. Summary LSAs appear in the routing table as IA (interarea routes).
- **E1** or **O E2** - External LSAs appear in the routing table marked as external type 1 (E1) or external type 2 (E2) routes.

OSPF Route Calculation

Each router uses the SPF algorithm against the LSDB to build the SPF tree. The SPF tree is used to determine the best paths.

The order in which the best paths are calculated is as follows:

1. All routers calculate the best paths to destinations within their area (intra-area) and add these entries to the routing table. These are the type 1 and type 2 LSAs, which are noted in the routing table with a routing designator of **O**.
2. All routers calculate the best paths to the other areas within the internetwork. These best paths are the interarea route entries, or type 3 and type 4 LSAs, and are noted with a routing designator of **O IA**.
3. All routers (except those that are in a form of stub area) calculate the best paths to the external autonomous system (type 5) destinations. These are noted with either an **O E1** or an **O E2** route designator, depending on the configuration.

When converged, a router can communicate with any network within or outside the OSPF autonomous system.

Topic 106: Configuring Multiarea OSPF:

Implementing Multiarea OSPF

OSPF can be implemented as single-area or multiarea. The type of OSPF implementation chosen depends on the specific requirements and existing topology.

There are 4 steps to implementing multiarea OSPF.

Steps 1 and 2 are part of the planning process.

Step 1. Gather the network requirements and parameters - This includes determining the number of host and network devices, the IP addressing scheme (if already implemented), the size of the routing domain, the size of the routing tables, the risk of topology changes, and other network characteristics.

Step 2. Define the OSPF parameters -Based on information gathered during Step 1, the network administrator must determine if single-area or multiarea OSPF is the preferred implementation. If multiarea OSPF is selected, there are several considerations the network administrator must take into account while determining the OSPF parameters, to include:

- **IP addressing plan** - This governs how OSPF can be deployed and how well the OSPF deployment might scale. A detailed IP addressing plan, along with the IP subnetting information, must be created. A good IP addressing plan should enable the usage of OSPF multiarea design and summarization. This plan more easily scales the network, as well as optimizes OSPF behavior and the propagation of LSA.
- **OSPF areas** - Dividing an OSPF network into areas decreases the LSDB size and limits the propagation of link-state updates when the topology changes. The routers that are to be ABRs and ASBRs must be identified, as are those that are to perform any summarization or redistribution.
- **Network topology** - This consists of links that connect the network equipment and belong to different OSPF areas in a multiarea OSPF design. Network topology is important to determine primary and backup links. Primary and backup links are defined by the changing OSPF cost on interfaces. A detailed network topology plan should also be used to determine the different OSPF areas, ABR, and ASBR as well as summarization and redistribution points, if multiarea OSPF is used.

Step 3. Configure the multiarea OSPF implementation based on the parameters.

Step 4. Verify the multiarea OSPF implementation based on the parameters.

Configuring Multiarea OSPF

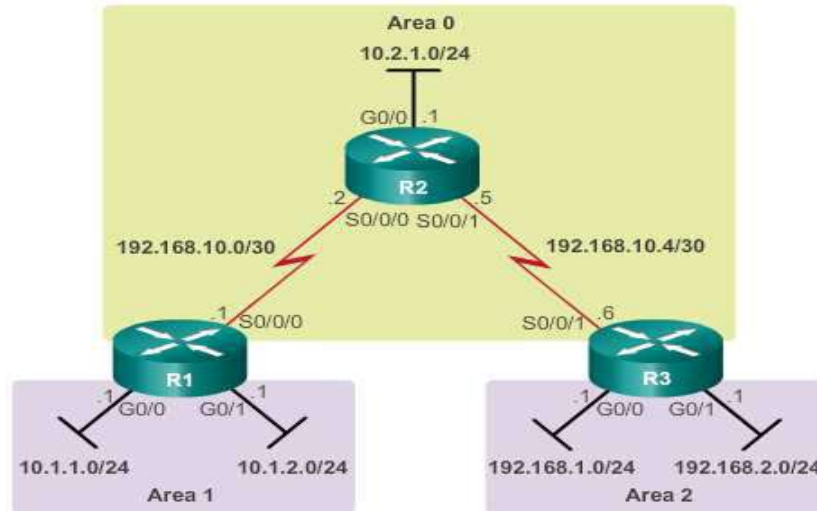


Figure above displays the reference multiarea OSPF topology. In this example:

- R1 is an ABR because it has interfaces in area 1 and an interface in area 0.
- R2 is an internal backbone router because all of its interfaces are in area 0.
- R3 is an ABR because it has interfaces in area 2 and an interface in area 0.

There are no special commands required to implement this multiarea OSPF network. A router simply becomes an ABR when it has two **network** statements in different areas.

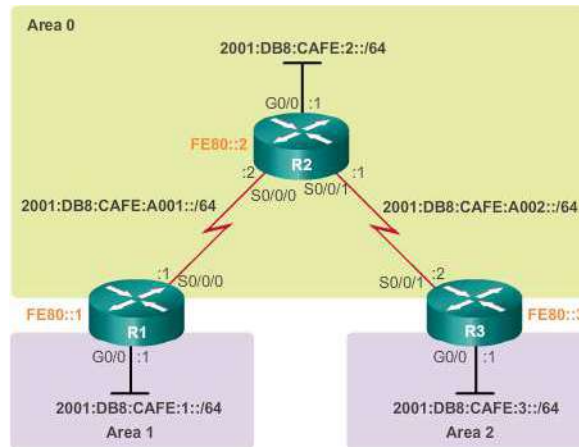
R1 is assigned the router ID 1.1.1.1. This example enables OSPF on the two LAN interfaces in area 1. The serial interface is configured as part of OSPF area 0. Because R2 has interfaces connected to two different areas, it is an ABR.

Upon completion of the R2 and R3 configuration, notice the informational messages informing of the adjacencies with R1 (1.1.1.1).

Upon completion of the R3 configuration, notice the informational messages informing of an adjacency with R1 (1.1.1.1) and R2 (2.2.2.2).

```
R1(config)# router ospf 10
R1(config-router)# router-id 1.1.1.1
R1(config-router)# network 10.1.1.1 0.0.0.0 area 1
R1(config-router)# network 10.1.2.1 0.0.0.0 area 1
R1(config-router)# network 192.168.10.1 0.0.0.0 area 0
R1(config-router)# end
R1#
```

Configuring Multiarea OSPFv3



Like OSPFv2, implementing the multiarea OSPFv3 topology in Figure above is simple. There are no special commands required. A router simply becomes an ABR when it has two interfaces in different areas.

R1 is assigned the router ID 1.1.1.1. The example also enables OSPF on the two LAN interfaces in area 1 and the serial interface in area 0. Because R1 has interfaces connected to two different areas, it becomes an ABR.

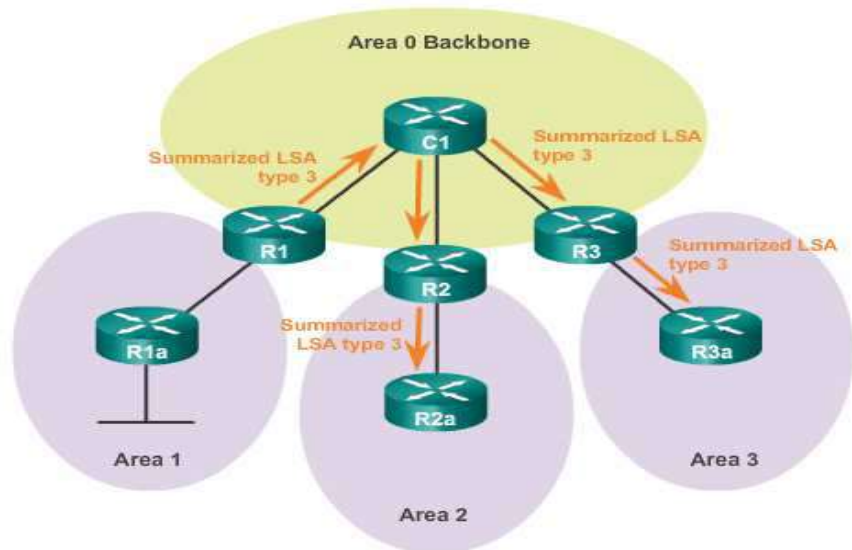
Upon completion of the R2 configuration, notice the message that there is an adjacency with R1 (1.1.1.1).

Upon completion of the R3 configuration, notice the message that there is an adjacency with R2 (2.2.2.2).

```
R1(config)# ipv6 router ospf 10
R1(config-rtr)# router-id 1.1.1.1
R1(config-rtr)# exit
R1(config)#
R1(config)# interface GigabitEthernet 0/0
R1(config-if)# ipv6 ospf 10 area 1
R1(config-if)#
R1(config-if)# interface Serial0/0/0
R1(config-if)# ipv6 ospf 10 area 0
R1(config-if)# end
R1#
```

Topic 107: OSPF Route Summarization:

Summarization helps keep routing tables small. It involves consolidating multiple routes into a single advertisement, which can then be propagated into the backbone area.

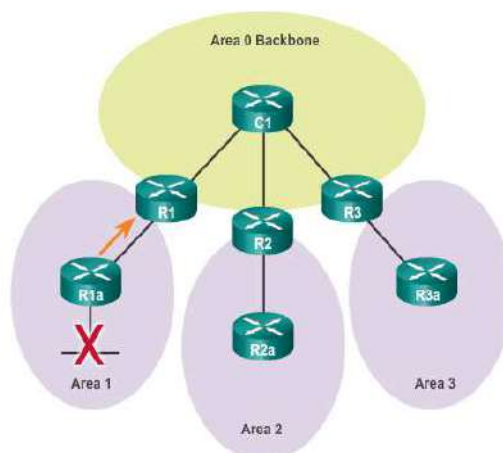


Normally, type 1 and type 2 LSAs are generated inside each area, translated into type 3 LSAs, and sent to other areas. If area 1 had 30 networks to advertise, then 30 type 3 LSAs would be forwarded into the backbone. With route summarization, the ABR consolidates the 30 networks into one of two advertisements.

In Figure above, R1 consolidates all of the network advertisements into one summary LSA. Instead of forwarding individual LSAs for each route in area 1, R1 forwards a summary LSA to the core router C1. C1 in turn, forwards the summary LSA to R2 and R3. R2 and R3 then forward it to their respective internal routers.

Summarization also helps increase the network's stability, because it reduces unnecessary LSA flooding. This directly affects the amount of bandwidth, CPU, and memory resources consumed by the OSPF routing process. Without route summarization, every specific-link LSA is propagated into the OSPF backbone and beyond, causing unnecessary network traffic and router overhead.

In Figure below, a network link on R1a fails. R1a sends an LSA to R1. However, R1 does not propagate the update, because it has a summary route configured. Specific-link LSA flooding outside the area does not occur.



Summarizing Interarea Routes on ABRs

In OSPF, summarization can only be configured on ABRs or ASBRs. Instead of advertising many specific networks, the ABR routers and ASBR routers advertise a summary route. ABR routers summarize type 3 LSAs and ASBR routers summarize type 5 LSAs.

By default, summary LSAs (type 3 LSAs) and external LSAs (type 5 LSAs) do not contain summarized (aggregated) routes; that is, by default, summary LSAs are not summarized.

Route summarization can be configured as follows:

Interarea route summarization - Interarea route summarization occurs on ABRs and applies to routes from within each area. It does not apply to external routes injected into OSPF via redistribution. To perform effective interarea route summarization, network addresses within areas should be assigned contiguously so that these addresses can be summarized into a minimal number of summary addresses.

External route summarization - External route summarization is specific to external routes that are injected into OSPF via route redistribution. Again, it is important to ensure the contiguity of the external address ranges that are being summarized. Generally, only ASBRs summarize external routes.

Note: External route summarization is configured on ASBRs using the **summary-address address mask** router configuration mode command.

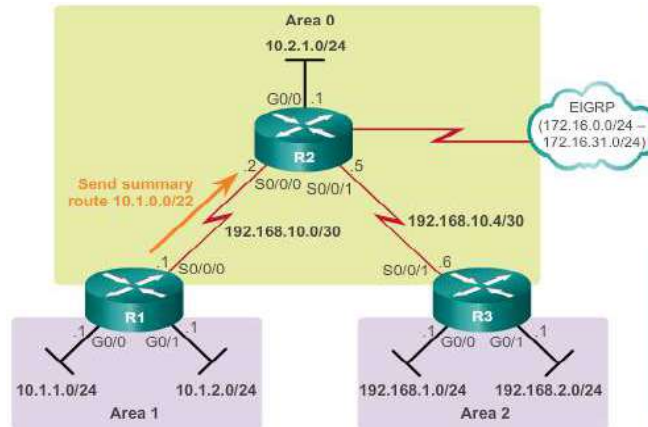
OSPF does not perform auto-summarization. Interarea summarization must be manually configured on ABRs.

Summarization of internal routes can only be done by ABRs. When summarization is enabled on an ABR, it injects into the backbone a single type 3 LSA describing the summary route. Multiple routes inside the area are summarized by the one LSA.

A summary route is generated if at least one subnet within the area falls in the summary address range. The summarized route metric is equal to the lowest cost of all subnets within the summary address range.

Note: An ABR can only summarize routes that are within the areas connected to the ABR.

Figure below shows a multiarea OSPF topology. The routing tables of R1 and R3 are examined to see the effect of the summarization.



```

R1# show ip route ospf | begin Gateway
Gateway of last resort is not set

    10.0.0.0/8 is variably subnetted, 5 subnets, 2 masks
O   10.2.1.0/24 [110/648] via 192.168.10.2, 00:00:49,
    Serial0/0/0
O IA 192.168.1.0/24 [110/1295] via 192.168.10.2, 00:00:49,
    Serial0/0/0
O IA 192.168.2.0/24 [110/1295] via 192.168.10.2, 00:00:49,
    Serial0/0/0
    192.168.10.0/24 is variably subnetted, 3 subnets, 2
    masks
O   192.168.10.4/30 [110/1294] via 192.168.10.2,
    00:00:49, Serial0/0/0
R1#

R3# show ip route ospf | begin Gateway
Gateway of last resort is not set

    10.0.0.0/24 is subnetted, 3 subnets
O IA 10.1.1.0 [110/1295] via 192.168.10.5, 00:27:14, Serial0/0/1
O IA 10.1.2.0 [110/1295] via 192.168.10.5, 00:27:14, Serial0/0/1
O   10.2.1.0 [110/648] via 192.168.10.5, 00:27:57, Serial0/0/1
    192.168.10.0/24 is variably subnetted, 3 subnets, 2 masks
O   192.168.10.0/30 [110/1294] via 192.168.10.5, 00:27:57,
    Serial0/0/1
R3#

```

Calculating the Summary Route

Step 1	Step 2	Some Bits Are Different
10.1.1.0	00001010.00000001.00000001.00000000	01.00000000
10.1.2.0	00001010.00000001.00000001.00000000	10.00000000

First 22 Bits Match

Summarized Route: 10.1.0.0/22 or 10.1.0.0 255.255.252.0

The figure above illustrates that summarizing networks into a single address and mask can be done in three steps:

Step 1. List the networks in binary format. In the example the two area 1 networks 10.1.1.0/24 and 10.1.2.0/24 are listed binary format.

Step 2. Count the number of far left matching bits to determine the mask for the summary route. As highlighted, the first 22 far left matching bits match. This results in the prefix **/22** or subnet mask **255.255.252.0**.

Step 3. Copy the matching bits and then add zero bits to determine the summarized network address. In this example, the matching bits with zeros at the end result in a network address of 10.1.0.0/22. This summary address summarizes four networks: 10.1.0.0/24, 10.1.1.0/24, 10.1.2.0/24, and 10.1.3.0/24.

In the example the summary address matches four networks although only two networks exist.

Summarizing Area 1 Routes on R1

```
R1(config)# router ospf 10
R1(config-router)# area 1 range 10.1.0.0 255.255.252.0
R1(config-router)#
```

Verify the R1 Routing Table After Summarization

```
R1# show ip route ospf | begin Gateway
Gateway of last resort is not set

    10.0.0.0/8 is variably subnetted, 6 subnets, 3 masks
O       10.1.0.0/22 is a summary, 00:00:09, Null0
O       10.2.1.0/24 [110/648] via 192.168.10.2, 00:00:09, Serial0/0/0
O IA    192.168.1.0/24 [110/1295] via 192.168.10.2, 00:00:09, Serial0/0/0
O IA    192.168.2.0/24 [110/1295] via 192.168.10.2, 00:00:09, Serial0/0/0
O       192.168.10.0/24 is variably subnetted, 3 subnets, 2 masks
O       192.168.10.4/30 [110/1294] via 192.168.10.2, 00:00:09, Serial0/0/0
R1#
```

Topic 108: Verifying Multiarea OSPF:

The same verification commands used to verify single-area OSPF also can be used to verify the multiarea OSPF topology:

show ip ospf neighbor

show ip ospf

show ip ospf interface

Commands that verify specific multiarea information include:

show ip protocols

show ip ospf interface brief

show ip route ospf

show ip ospf database

Note: For the equivalent OSPFv3 command, simply substitute **ip** with **ipv6**.

Use the **show ip protocols** command to verify the OSPF status. The output of the command reveals which routing protocols are configured on a router. It also includes routing protocol specifics such as the

router ID, number of areas in the router, and networks included within the routing protocol configuration.

Use the **show ip ospf interface brief** command to display concise OSPF-related information of OSPF-enabled interfaces. This command reveals useful information, such as the OSPF process ID that the interface is assigned to, the area that the interfaces are in, and the cost of the interface.

Verify the OSPF Routes

The most common command used to verify a multiarea OSPF configuration is the **show ip route** command. Add the **ospf** parameter to display only OSPF-related information.

Use the **show ip ospf database** command to verify the contents of the LSDB.

There are many command options available with the **show ip ospf database** command.

Verify Multiarea OSPFv3

Like OSPFv2, OSPFv3 provides similar OSPFv3 verification commands.

Topic 109: Packet Tracer – Configuring Multiarea OSPFv2:

- Configure Multiarea OSPFv2 Verify Configuration

Topic 110: Packet Tracer – Configuring Multiarea OSPFv2 - 2:

This is the continuation of the previous topic.

Topic 111: Packet Tracer – Configuring Multiarea OSPFv2 - 3:

This is the continuation of the previous topic.

Topic 112: IP ACL:

Access Control Lists (ACLs)

An ACL is a series of IOS commands that control whether a router forwards or drops packets based on information found in the packet header. ACLs are among the most commonly used features of Cisco IOS software.

When configured, ACLs perform the following tasks:

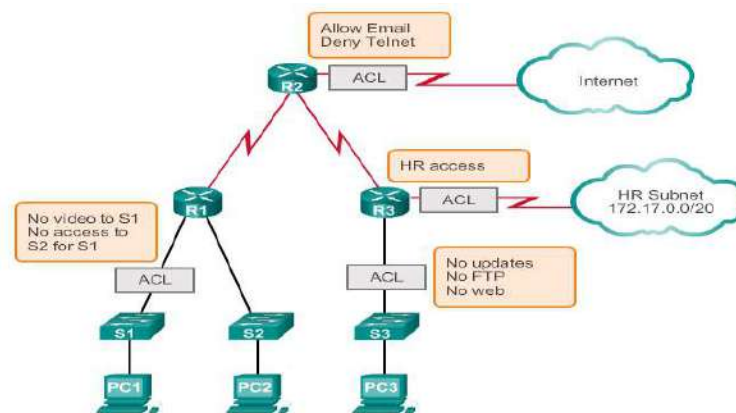
- Limit network traffic to increase network performance. For example, if corporate policy does not allow video traffic on the network, ACLs that block video traffic could be configured and applied. This would greatly reduce the network load and increase network performance.
- Provide traffic flow control. ACLs can restrict the delivery of routing updates. If updates are not required because of network conditions, bandwidth is preserved.
- Provide a basic level of security for network access. ACLs can allow one host to access a part of the network and prevent another host from accessing the same area. For example, access to the Human Resources network can be restricted to authorized users.
- Filter traffic based on traffic type. For example, an ACL can permit email traffic, but block all Telnet traffic.

- Screen hosts to permit or deny access to network services. ACLs can permit or deny a user to access file types, such as FTP or HTTP.

By default, a router does not have ACLs configured; therefore, by default a router does not filter traffic. Traffic that enters the router is routed solely based on information within the routing table. However, when an ACL is applied to an interface, the router performs the additional task of evaluating all network packets as they pass through the interface to determine if the packet can be forwarded.

In addition to either permitting or denying traffic, ACLs can be used for selecting types of traffic to be analyzed, forwarded, or processed in other ways. For example, ACLs can be used to classify traffic to enable priority processing. This capability is similar to having a VIP pass at a concert or sporting event. The VIP pass gives selected guests privileges not offered to general admission ticket holders, such as priority entry or being able to enter a restricted area.

The figure below shows a sample topology with ACLs applied.



Packet Filtering

So how does an ACL use the information passed during a TCP/IP conversation to filter traffic?

Packet filtering, sometimes called static packet filtering, controls access to a network by analyzing the incoming and outgoing packets and passing or dropping them based on given criteria, such as the source IP address, destination IP addresses, and the protocol carried within the packet.

A router acts as a packet filter when it forwards or denies packets according to filtering rules. When a packet arrives at the packet-filtering router, the router extracts certain information from the packet header. Using this information, the router makes decisions, based on configured filter rules, as to whether the packet can pass through or be discarded.

A packet-filtering router uses rules to determine whether to permit or deny traffic. A router can also perform packet filtering at Layer 4, the transport layer. The router can filter packets based on the source port and destination port of the TCP or UDP segment. These rules are defined using ACLs.

An ACL is a sequential list of permit or deny statements, known as access control entries (ACEs). ACEs are also commonly called ACL statements. ACEs can be created to filter traffic based on certain criteria such as: the source address, destination address, the protocol, and port numbers. When network traffic

passes through an interface configured with an ACL, the router compares the information within the packet against each ACE, in sequential order, to determine if the packet matches one of the statements. If a match is found, the packet is processed accordingly. In this way, ACLs can be configured to control access to a network or subnet.

To evaluate network traffic, the ACL extracts the following information from the Layer 3 packet header:

- Source IP address
- Destination IP address
- ICMP message type

The ACL can also extract upper layer information from the Layer 4 header, including:

- TCP/UDP source port
- TCP/UDP destination port
- ACL Operation

ACLs define the set of rules that give added control for packets that enter inbound interfaces, packets that relay through the router, and packets that exit outbound interfaces of the router. ACLs do not act on packets that originate from the router itself.

ACLs are configured to apply to inbound traffic or to apply to outbound traffic.

Inbound ACLs - Incoming packets are processed before they are routed to the outbound interface. An inbound ACL is efficient because it saves the overhead of routing lookups if the packet is discarded. If the packet is permitted by the tests, it is then processed for routing. Inbound ACLs are best used to filter packets when the network attached to an inbound interface is the only source of the packets needed to be examined.

Outbound ACLs - Incoming packets are routed to the outbound interface, and then they are processed through the outbound ACL. Outbound ACLs are best used when the same filter will be applied to packets coming from multiple inbound interfaces before exiting the same outbound interface.

The last statement of an ACL is always an implicit deny. This statement is automatically inserted at the end of each ACL even though it is not physically present. The implicit deny blocks all traffic. Because of this implicit deny, an ACL that does not have at least one permit statement will block all traffic.

Topic 113: Standard vs. Extended IPv4 ACL:

The two types of Cisco IPv4 ACLs are standard and extended.

Standard ACLs

Standard ACLs can be used to permit or deny traffic only from source IPv4 addresses. The destination of the packet and the ports involved are not evaluated.

Extended ACLs

Extended ACLs filter IPv4 packets based on several attributes:

- Protocol type

- Source IPv4 address
- Destination IPv4 address
- Source TCP or UDP ports
- Destination TCP or UDP ports
- Optional protocol type information for finer control

Numbered and Named ACLs

Standard and extended ACLs can be created using either a number or a name to identify the ACL and its list of statements.

Using numbered ACLs is an effective method for determining the ACL type on smaller networks with more homogeneously defined traffic. However, a number does not provide information about the purpose of the ACL.

Wildcard Masking

IPv4 ACEs include the use of wildcard masks. A wildcard mask is a string of 32 binary digits used by the router to determine which bits of the address to examine for a match.

As with subnet masks, the numbers 1 and 0 in the wildcard mask identify how to treat the corresponding IP address bits. However, in a wildcard mask, these bits are used for different purposes and follow different rules.

- Subnet masks use binary 1s and 0s to identify the network, subnet, and host portion of an IP address. Wildcard masks use binary 1s and 0s to filter individual IP addresses or groups of IP addresses to permit or deny access to resources.
- Wildcard masks and subnet masks differ in the way they match binary 1s and 0s. Wildcard masks use the following rules to match binary 1s and 0s:
 - Wildcard mask bit 0 - Match the corresponding bit value in the address.
 - Wildcard mask bit 1 - Ignore the corresponding bit value in the address.

Wildcard Bit Mask Keywords

Working with decimal representations of binary wildcard mask bits can be tedious. To simplify this task, the keywords **host** and **any** help identify the most common uses of wildcard masking. These keywords eliminate entering wildcard masks when identifying a specific host or an entire network. These keywords also make it easier to read an ACL by providing visual clues as to the source or destination of the criteria.

The **host** keyword substitutes for the 0.0.0.0 mask. This mask states that all IPv4 address bits must match or only one host is matched.

The **any** option substitutes for the IP address and 255.255.255.255 mask. This mask says to ignore the entire IPv4 address or to accept any addresses.

Topic 114: ACL Creation and Placement:

Writing ACLs can be a complex task. For every interface there may be multiple policies needed to manage the type of traffic allowed to enter or exit that interface.

Here are some guidelines for using ACLs:

- Use ACLs in firewall routers positioned between your internal network and an external network such as the Internet.
- Use ACLs on a router positioned between two parts of your network to control traffic entering or exiting a specific part of your internal network.
- Configure ACLs on border routers, that is, routers situated at the edges of your networks. This provides a very basic buffer from the outside network, or between a less controlled area of your own network and a more sensitive area of your network.
- Configure ACLs for each network protocol configured on the border router interfaces.

The Three Ps

A general rule for applying ACLs on a router can be recalled by remembering the three Ps. You can configure one ACL per protocol, per direction, per interface:

One ACL per protocol -To control traffic flow on an interface, an ACL must be defined for each protocol enabled on the interface.

One ACL per direction -ACLs control traffic in one direction at a time on an interface. Two separate ACLs must be created to control inbound and outbound traffic.

One ACL per interface - ACLs control traffic for an interface, for example, GigabitEthernet 0/0.

Where to Place ACLs

The proper placement of an ACL can make the network operate more efficiently. An ACL can be placed to reduce unnecessary traffic. For example, traffic that will be denied at a remote destination should not be forwarded using network resources along the route to that destination.

Every ACL should be placed where it has the greatest impact on efficiency. The basic rules are:

Extended ACLs - Locate extended ACLs as close as possible to the source of the traffic to be filtered. This way, undesirable traffic is denied close to the source network without crossing the network infrastructure.

Standard ACLs - Because standard ACLs do not specify destination addresses, place them as close to the destination as possible. Placing a standard ACL at the source of the traffic will effectively prevent that traffic from reaching any other networks through the interface where the ACL is applied.

Placement of the ACL and therefore the type of ACL used may also depend on:

The extent of the network administrator's control - Placement of the ACL can depend on whether or not the network administrator has control of both the source and destination networks.

Bandwidth of the networks involved- Filtering unwanted traffic at the source prevents transmission of the traffic before it consumes bandwidth on the path to a destination. This is especially important in low bandwidth networks.

Ease of configuration - If a network administrator wants to deny traffic coming from several networks, one option is to use a single standard ACL on the router closest to the destination. The disadvantage is that traffic from these networks will use bandwidth unnecessarily. An extended ACL could be used on each router where the traffic originated. This will save bandwidth by filtering the traffic at the source but requires creating extended ACLs on multiple routers.

Standard ACL Placement

A standard ACL can only filter traffic based on a source address. The basic rule for placement of a standard ACL is to place the ACL as close as possible to the destination network. This allows the traffic to reach all other networks except the network where the packets will be filtered.

Extended ACL Placement

Like a standard ACL, an extended ACL can filter traffic based on the source address. However, an extended ACL can also filter traffic based on the destination address, protocol, and port number. This allows network administrators more flexibility in the type of traffic that can be filtered and where to place the ACL. The basic rule for placing an extended ACL is to place it as close to the source as possible. This prevents unwanted traffic from being sent across multiple networks only to be denied when it reaches its destination. Network administrators can only place ACLs on devices that they control. Therefore, placement must be determined in the context of where the control of the network administrator extends.

Topic 115: Configure Standard IPv4 ACLs:

How are ACLs Created?

In Two Steps

1. Create an ACL definition.

- Enter global configuration mode.
- Define statements of what to filter.

2. Apply the ACL to an interface.

- Enter interface configuration mode.
- Identify the ACL and the direction to filter.

Create a Standard ACL

```
RTR(config)# access-list ACL# {permit|deny} { test-conditions }
```

```
access-list 5 permit 172.34.54.34 0.0.0.0
```

Apply the ACL to an Interface

RTR(config-if)# {protocol}	access-group	list-#	{in out}
ip	access-group	5	out

Implied Deny

When traffic enters the router, the traffic is compared to all ACEs in the order that the entries occur in the ACL. The router continues to process the ACEs until it finds a match. The router will process the packet based on the first match found and no other ACEs will be examined.

If no matches are found when the router reaches the end of the list, the traffic is denied. This is because, by default, there is an implied deny at the end of all ACLs for traffic that was not matched to a configured entry. A single-entry ACL with only one deny entry has the effect of denying all traffic. At least one permit ACE must be configured in an ACL or all traffic is blocked.

Configuring Standard ACLs

To use numbered standard ACLs on a Cisco router, you must first create the standard ACL and then activate the ACL on an interface.

The **access-list** global configuration command defines a standard ACL with a number in the range of 1 through 99. Cisco IOS Software Release 12.0.1 extended these numbers by allowing 1300 to 1999 to be used for standard ACLs. This allows for a maximum of 798 possible standard ACLs. These additional numbers are referred to as expanded IP ACLs.

The full syntax of the standard ACL command is as follows:

```
Router(config)# access-list access-list-number { deny | permit | remark } source [source-wildcard] [log ]
```

ACEs can deny or permit an individual host or a range of host addresses. To create a host statement in numbered ACL 10 that permits a specific host with the IP address 192.168.10.0, you would enter:

```
R1(config)# access-list 10 permit host 192.168.10.10
```

To create a statement that will permit a range of IPv4 addresses in a numbered ACL 10 that permits all IPv4 addresses in the network 192.168.10.0/24, you would enter:

```
R1(config)# access-list 10 permit 192.168.10.0 0.0.0.255
```

To remove the ACL, the global configuration **no access-list** command is used. Issuing the **show access-list** command confirms that access list 10 has been removed.

Typically, when an administrator creates an ACL, the purpose of each statement is known and understood. However, to ensure that the administrator and others recall the purpose of a statement,

remarks should be included. The **remark** keyword is used for documentation and makes access lists a great deal easier to understand. Each remark is limited to 100 characters.

Internal Logic – Order Matters

Cisco IOS applies an internal logic when accepting and processing standard ACEs. ACEs are processed sequentially. Therefore, the order in which ACEs are entered is important.

Topic 116: Configure Standard IPv4 ACLs - 2:

Applying Standard ACLs to Interfaces

Standard ACL Configuration Procedures

After a standard ACL is configured, it is linked to an interface using the **ip access-group** command in interface configuration mode:

```
Router(config-if)# ip access-group { access-list-number | access-list-name } { in | out }
```

To remove an ACL from an interface, first enter the **no ip access-group** command on the interface, and then enter the global **no access-list** command to remove the entire ACL.

Standard ACL Configuration Procedures

After a standard ACL is configured, it is linked to an interface using the **ip access-group** command in interface configuration mode:

```
Router(config-if)# ip access-group { access-list-number | access-list-name } { in | out }
```

To remove an ACL from an interface, first enter the **no ip access-group** command on the interface, and then enter the global **no access-list** command to remove the entire ACL.

Creating a Named ACL

Naming an ACL makes it easier to understand its function. For example, an ACL configured to deny FTP could be called NO_FTP. When you identify your ACL with a name instead of with a number, the configuration mode and command syntax are slightly different.

The steps required to create a standard named ACL.

Step 1. Starting from the global configuration mode, use the **ip access-list** command to create a named ACL. ACL names are alphanumeric, case sensitive, and must be unique. The **ip access-list standard name** is used to create a standard named ACL, whereas the command **ip access-list extended name** is for an extended access list. After entering the command, the router is in named standard ACL configuration mode as indicated by the prompt.

Note: Numbered ACLs use the global configuration command **access-list** whereas named IPv4 ACLs use the **ip access-list** command.

Step 2. From the named ACL configuration mode, use **permit** or **deny** statements to specify one or more conditions for determining whether a packet is forwarded or dropped.

Step 3. Apply the ACL to an interface using the **ip access-group** command. Specify if the ACL should be applied to packets as they enter into the interface (**in**) or applied to packets as they exit the interface (**out**).

Commenting ACLs

You can use the **remark** keyword to include comments (remarks) about entries in any IP standard or extended ACL. The remarks make the ACL easier for you to understand and scan. Each remark line is limited to 100 characters.

The remark can go before or after a **permit** or **deny** statement. You should be consistent about where you put the remark so that it is clear which remark describes which **permit** or **deny** statement. For example, it would be confusing to have some remarks before the associated **permit** or **deny** statements and some remarks after the statements.

To include a comment for IPv4 numbered standard or extended ACLs, use the **access-list access-list_number remark remark** global configuration command. To remove the remark, use the **no** form of this command.

Topic 117: Packet Tracer – Configuring Standard ACLs:

- Plan an ACL Implementation
- Configure, Apply and Verify a Standard ACL

Topic 118: Modify IPv4 ACLs:

Editing Numbered ACLs

When configuring a standard ACL, the statements are added to the running-config. However, there is no built-in editing feature that allows you to edit a change in an ACL.

There are two ways that a standard numbered ACL can be edited.

Method 1: Using a Text Editor

After someone is familiar with creating and editing ACLs, it may be easier to construct the ACL using a text editor such as Microsoft Notepad. This allows you to create or edit the ACL and then paste it into the router. For an existing ACL, you can use the **show running-config** command to display the ACL, copy and paste it into the text editor, make the necessary changes, and paste it back in.

Configuration: Here are the steps to edit and correct ACL 1:

Step 1. Display the ACL using the **show running-config** command.

Step 2. Highlight the ACL, copy it, and then paste it into Microsoft Notepad. Edit the list as required. After the ACL is correctly displayed in Microsoft Notepad, highlight it and copy it.

Step 3. In global configuration mode, remove the access list using the **no access-list 1** command. Otherwise, the new statements would be appended to the existing ACL. Then paste the new ACL into the configuration of the router.

Step 4. Using the **show running-config** command, verify the changes

Method 2: Using the Sequence Number

The initial configuration of ACL 1 included a host statement for host 192.168.10.99. This was in error. The host should have been configured as 192.168.10.10. To edit the ACL using sequence numbers follow these steps:

Step 1. Display the current ACL using the **show access-lists 1** command. The output from this command will be discussed in more detail later in this section. The sequence number is displayed at the beginning of each statement. The sequence number was automatically assigned when the access list statement was entered. Notice that the misconfigured statement has the sequence number 10.

Step 2. Enter the **ip access-lists standard** command that is used to configure named ACLs. The ACL number, 1, is used as the name. First the misconfigured statement needs to be deleted using the **no 10** command with 10 referring to the sequence number. Next, a new sequence number 10 statement is added using the command, **10 deny host 192.168.10.10**.

Note: Statements cannot be overwritten using the same sequence number as an existing statement. The current statement must be deleted first, and then the new one can be added.

Step 3. Verify the changes using the **show access-lists** command.

Verifying ACLs

The **show ip interface** command is used to verify the ACL on the interface. The output from this command includes the number or name of the access list and the direction in which the ACL was applied.

Viewing ACLs Statistics

Once the ACL has been applied to an interface and some testing has occurred, the **show access-lists** command will show statistics for each statement that has been matched.

During testing of an ACL, the counters can be cleared using the **clear access-list counters** command. This command can be used alone or with the number or name of a specific ACL.

Topic 119: Extended IPv4 ACLs:

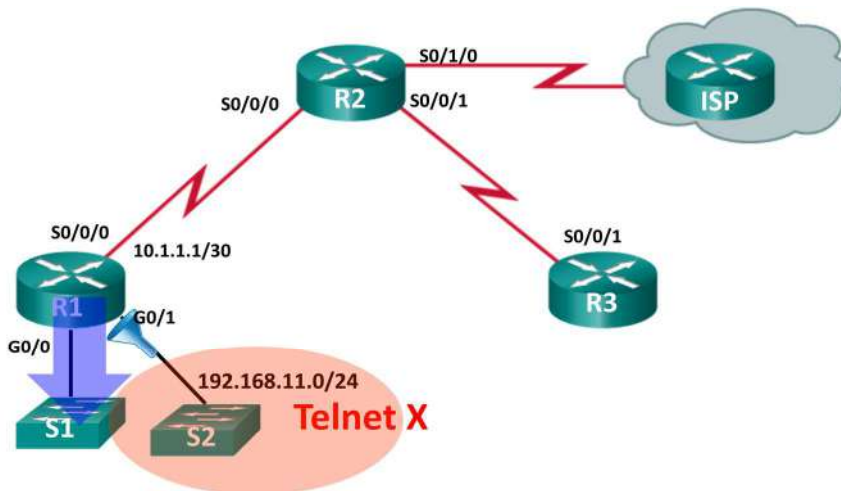
Testing Packets with Extended ACLs

For more precise traffic-filtering control, extended IPv4 ACLs can be created. Extended ACLs are numbered 100 to 199 and 2000 to 2699, providing a total of 799 possible extended numbered ACLs. Extended ACLs can also be named.

Extended ACLs are used more often than standard ACLs because they provide a greater degree of control.

Topic 120: Extended IPv4 ACLs - 2:

Deny Telnet and Permit Everything Else



```

R1(config)# access-list 102 deny tcp any 192.168.11.0 0.0.0.255 eq 23
R1(config)# access-list 102 permit ip any any

R1(config)# interface g0/0
R1(config-if)# ip access-group 102 out

```

In this example, the network administrator configured an ACL to allow users from the 192.168.10.0/24 network to browse both insecure and secure websites. Even though it has been configured, the ACL will not filter traffic until it is applied to an interface. To apply an ACL to an interface, first consider whether the traffic to be filtered is going in or out. When a user on the internal LAN accesses a website on the Internet, traffic is traffic going out to the Internet. When an internal user receives an email from the Internet, traffic is coming into the local router. However, when applying an ACL to an interface, in and out take on different meanings. From an ACL consideration, in and out are in reference to the router interface.

In the topology in the figure above, R1 has three interfaces. It has a serial interface, S0/0/0, and two Gigabit Ethernet interfaces, G0/0 and G0/1. Recall that an extended ACL should typically be applied close to the source. In this topology the interface closest to the source of the target traffic is the G0/0 interface.

Web request traffic from users on the 192.168.10.0/24 LAN is inbound to the G0/0 interface. Return traffic from established connections to users on the LAN is outbound from the G0/0 interface. The example applies the ACL to the G0/0 interface in both directions. The inbound ACL, 103, checks for the type of traffic. The outbound ACL, 104, checks for return traffic from established connections. This will restrict 192.168.10.0 Internet access to allow only website browsing.

Creating Named Extended ACLs

Named extended ACLs are created in essentially the same way that named standard ACLs are created. Follow these steps to create an extended ACL, using names:

Step 1. From global configuration mode, use the **ip access-list extended name** command to define a name for the extended ACL.

Step 2. In named ACL configuration mode, specify the conditions to **permit** or **deny**.

Step 3. Return to privileged EXEC mode and verify the ACL with the **show access-lists name** command.

Step 4. Save the entries in the configuration file with the **copy running-config startup-config** command.

To remove a named extended ACL, use the **no ip access-list extended name** global configuration command.

Verifying Extended ACLs

After an ACL has been configured and applied to an interface, use Cisco IOS **show** commands to verify the configuration.

Unlike standard ACLs, extended ACLs do not implement the same internal logic and hashing function. The output and sequence numbers displayed in the **show access-lists** command output is the order in which the statements were entered. Host entries are not automatically listed prior to range entries.

After an ACL configuration has been verified, the next step is to confirm that the ACLs work as planned; blocking and permitting traffic as expected.

Topic 121: Packet Tracer – Configuring Extended ACLs:

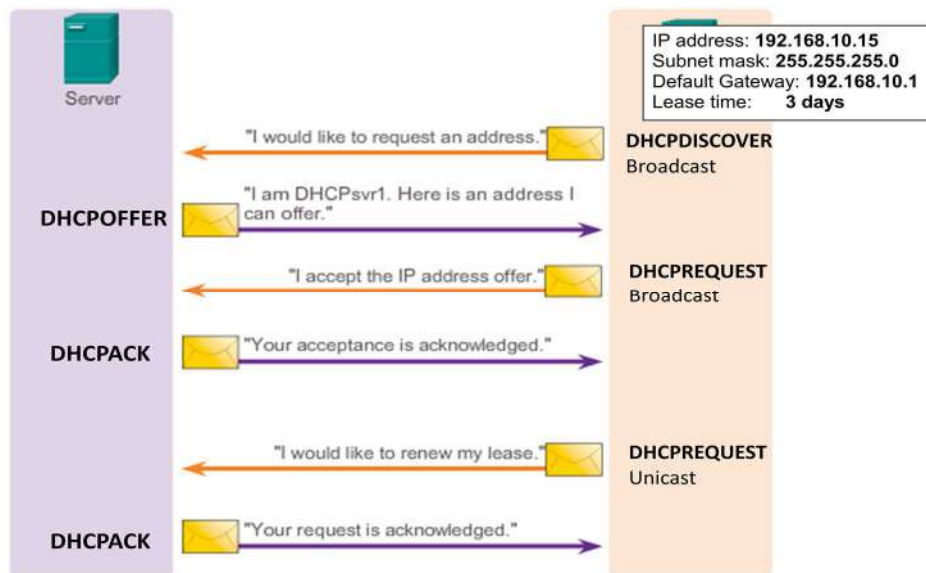
- Configure, Apply, and Verify an Extended Numbered ACL

Topic 122: Packet Tracer – Configuring Extended ACLs - 2:

This is the continuation of the previous topic.

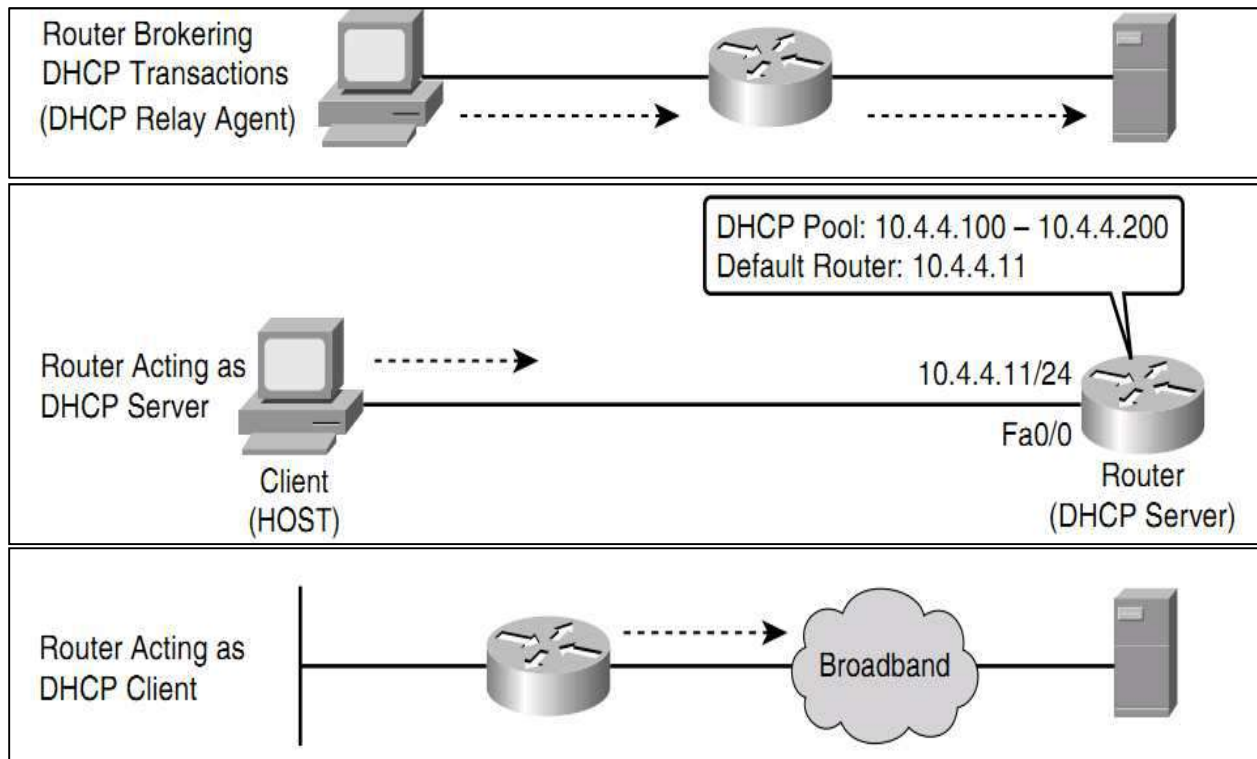
Topic 123: Dynamic Host Configuration Protocol (DHCPv4):

- Hosts were assigned an IP address manually.
- Next came BOOTP.
- The Bootstrap Protocol (BOOTP), defined in RFC 951, is the predecessor of DHCP.
- BOOTP provided a method for diskless workstations to download IP address configurations.
 - Diskless workstations are workstation or PCs without disk drives
 - E.g., Thin clients, cash register systems.
- Now we have DHCP.
- DHCP emerged as a standard protocol in October 1993 as defined in RFC 1531, succeeding the BOOTP protocol.
- DHCP allows a host to quickly and dynamically obtain configuration parameters from a DHCP Server.
 - In an Enterprise, DHCP is almost exclusively always configured on a dedicated server.



Topic 124: Configuring DHCPv4:

Router DHCP Roles



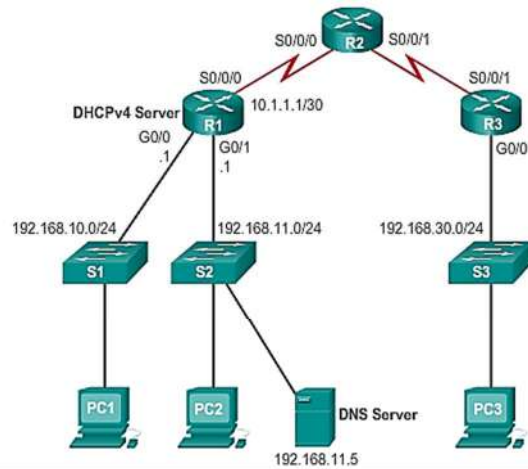
IOS DHCP Server

- The DHCP service is enabled by default on newer IOS.
- To disable DHCP, in global config mode:
 - **no service dhcp**

- To re-enable:
 - **service dhcp global**

Steps to Configuring DHCP on a Router

1. Define the excluded address ranges.



```
R1(config)#ip dhcp excluded-address low-address [high-address]
```

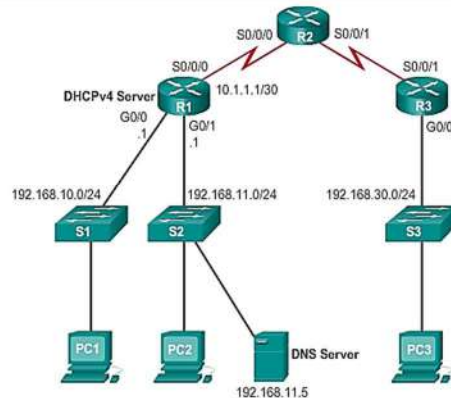
```
R1(config)# ip dhcp excluded-address 192.168.10.1 192.168.10.9
```

```
R1(config)# ip dhcp excluded-address 192.168.10.254
```

2. Create the DHCP address pools.

- Create the DHCP pool using the **ip dhcp pool** command.
- Enters you into DHCP configuration mode

```
R1(config)#ip dhcp pool POOL-NAME
```



```
R1(config)#ip dhcp pool LAN-POOL-1
```

```
R1(dhcp-config)#
```

3. Configure the specifics of the pool.

- Enter DHCP configuration mode and configure the specifics

Tasks	Command
Define the address pool	network <i>network-number</i> [<i>mask</i> <i>lprefix-length</i>]
Define the default gateway.	default-router <i>address</i> [<i>address2...address8</i>]

Optional Tasks	Command
Define a DNS server.	dns-server <i>address</i> [<i>address2...address8</i>]
Define the domain name.	domain-name <i>domain</i>
Define the duration of the DHCP lease.	lease { <i>days</i> [<i>hours</i>] [<i>minutes</i>] <i>infinite</i> }
Define an option code	option <i>code</i> [<i>instance number</i>] { <i>ascii string</i> <i>hex string</i> <i>ip-address</i> } option 150 <i>ip</i> 192.168.1.254

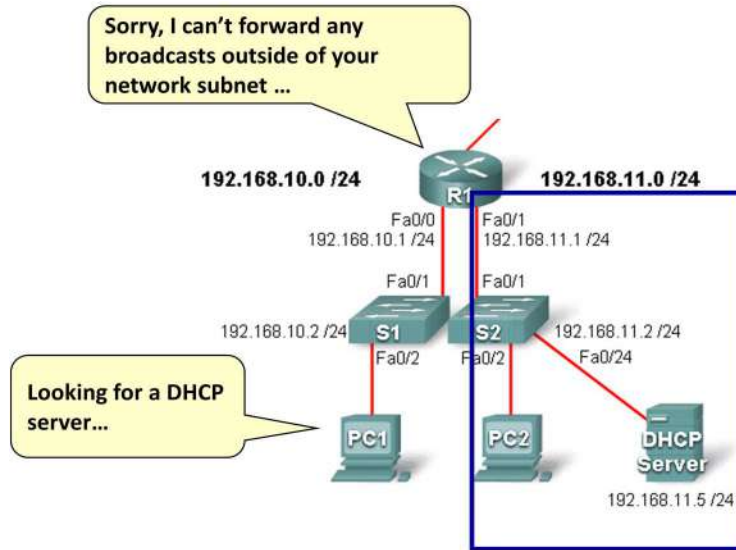
Topic 125: Configuring DHCPv4 - 2:

Verify DHCP Configuration

- **show ip dhcp server statistics**
 - Displays counts for server statistics and messages sent and received for an IOS-based DHCP server.
- **show ip dhcp binding**
 - Displays DHCP binding information for IP address assignment and subnet allocation.
- **show ip dhcp conflict**
 - Displays address conflicts found by a Cisco IOS DHCP server when addresses are offered to the client.
- **show ip dhcp pool *name***
 - Displays the subnets pool allocated and the current utilization level for the pool or all the pools if the name argument is not used.

Topic 126: DHCPv4 Relay:

DHCP Problems



DHCP Relays – Helper Addresses

Remote clients may require DHCP services and send broadcasts to locate these servers.

Routers, by default, will not forward client broadcasts beyond their subnet.

Solution:

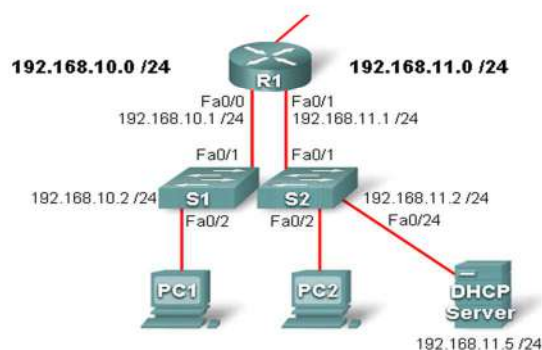
- Place DHCP servers on all subnets
- Use the Cisco IOS helper address feature.
- The IP helper address enables a router forward a UDP broadcast to a specific unicast IP address.
- Configured using the **ip helper-address** interface configuration command.

Command relays UDP broadcast requests.

Configured on the interface receiving the broadcast.

DHCP Relays

```
R1# config t
R1(config)# interface Fa0/0
R1(config-if)# ip helper-address 192.168.11.5
R1(config-if)# end
```



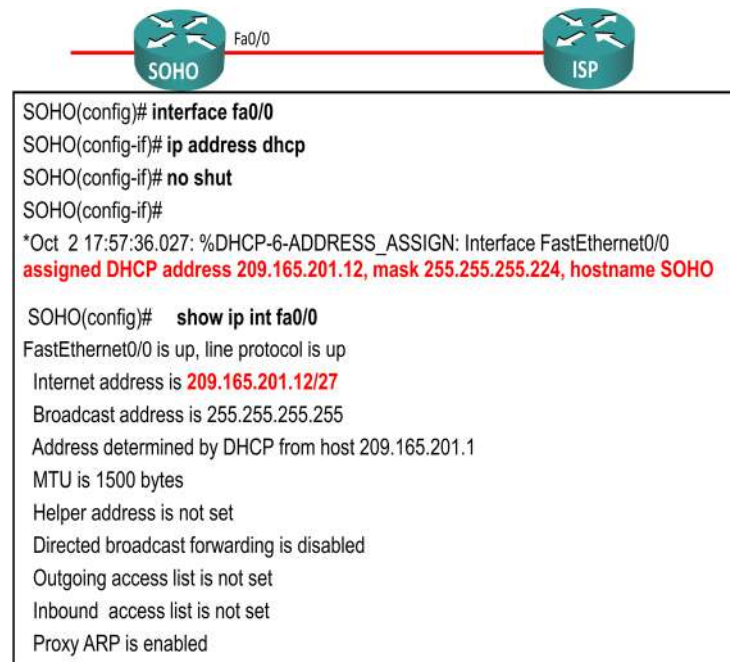
Topic 127: Configure DHCPv4 Client:

- Routers can also be DHCP clients.
 - Typical in small broadband routers for home use to connect to an ISP using a DSL or cable modem.
 - They are set to acquire an IP address automatically from their ISPs.

Router(config-if)#

```
ip address dhcp
```

Enables a Cisco IOS device to obtain an IP address dynamically from a DHCP server



Topic 128: Packet Tracer – Configure Basic DHCPv4 on a Router-1:

- Build a Network
- Configure DHCPv4 Server and Relay Agent

Topic 129: Packet Tracer – Configure Basic DHCPv4 on a Router-2:

This is the continuation of the previous topic.

Topic 130: Packet Tracer – Configure Basic DHCPv4 on a Router-3:

This is the continuation of the previous topic.

Topic 131: NAT

Internet Concerns

- There are not enough public IPv4 addresses to assign a unique address to each device connected to the Internet.
- Therefore the IETF developed several solutions to help stave off this depletion of global IPv4 addresses:
 - Subnetting
 - Variable-length subnet masking (VLSM)
 - Classless interdomain routing (CIDR)
 - Route summarization
 - Private addressing and NAT
 - Long term solution: IP version 6 (IPv6)

Private Addresses

Class	RFC 1918 Range		CIDR Prefix
A	10.0.0.0	– 10.255.255.255	10.0.0.0/8
B	172.16.0.0	– 172.31.255.255	172.16.0.0/12
C	192.168.0.0	– 192.168.255.255	192.168.0.0/16

Network Address Translation (NAT)

- NAT is a process used to translate network addresses.
- NAT's primary use is to conserve public IPv4 addresses.
- NAT is usually implemented on border network devices (e.g., firewalls or routers).
- NAT allows the networks to use private addresses internally, only translating to public addresses when needed.
- Devices within the organization are assigned private addresses and operate with locally unique addresses.
- When traffic must be sent or received to or from other organizations or the Internet, the border router translates the addresses to a public and globally unique address.
- You need to connect to the Internet and your hosts do not have globally unique IP addresses.
- You change over to a new Internet service provider (ISP) that requires you to renumber your network.
- Two intranets with duplicate addresses merge.
- You want to support basic load sharing.

NAT Terminology

NAT terminology is always applied from the perspective of the device with the translated address:

Inside address: The address of the device which is being translated by NAT.

Outside address: The address of the destination device.

NAT also uses the concept of local or global with respect to addresses:

Local address: A local address is any address that appears on the inside portion of the network.

Global address: A global address is any address that appears on the outside portion of the network.

Topic 132: Types of NAT Applications – Static NAT:

NAT Application Types

- **Static address translation (static NAT):**
 - One-to-one address mapping between local and global addresses.
- **Dynamic address translation (dynamic NAT):**
 - Many-to-many address mapping between local and global addresses.
- **Port Address Translation (PAT):**
 - Many-to-one address mapping between local and global addresses.
 - This method is also known as overloading (NAT overloading).

Static NAT

- Permanently bind an inside local address to an inside global address.
- Mappings are configured by the administrator and remain constant.
- Typically used to configure an internal server that must be accessed from the outside world.

Static NAT Configuration Steps

1. Create a mapping between the inside local address and the inside global addresses.
 - Configured using the **ip nat inside source static** *inside-local inside-global* global configuration command.
2. Identify the inside and outside NAT interfaces.
 - Configured using the **ip nat inside** and **ip nat outside** interface configuration commands.

Verifying Static NAT Example

```
R2# show ip nat translations
Pro Inside global    Inside local    Outside local    Outside global
--- 209.165.201.5    192.168.10.254    ---             ---
R2#
```

Topic 133: Types of NAT Applications - Dynamic NAT:

Dynamic NAT uses a pool of public addresses and assigns them on a first-come, first-served basis.

When an inside device requests access to an outside network, dynamic NAT assigns the inside local address an inside global address from a pool of addresses.

Dynamic NAT Configuration Steps

1. Define the pool of addresses that will be used for translation.

- Configured using the **ip nat pool** *pool-name start-ip end-ip {netmask netmask / prefix-length prefix-length}* global configuration command.
- 2. Configure a standard ACL to identify (permit) only those addresses that are to be translated.
- 3. Bind the ACL to the pool.
 - Configured using the **ip nat inside source list** *acl-# pool pool-name* global config command.
- 4. Identify the inside and outside NAT interfaces.
 - Configured using the **ip nat inside** and **ip nat outside** interface configuration commands.

Dynamic NAT Timeout

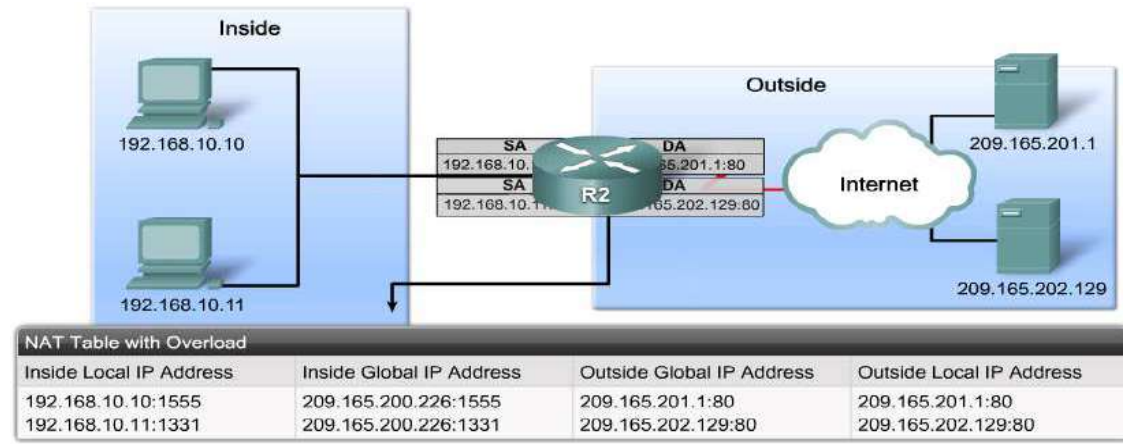
- Dynamic translations are temporary, and will eventually time out (default 24 hours).
 - Timeout can be configured.
 - It is important for translation table entries to time out so that addresses in the pool become available for other hosts.
 - If translation table entries do not time out fast enough, the entire pool of addresses could be in use.

```
Router(config)# ip nat translation timeout sec
Router(config)# ip nat translation timeout 120
```

Topic 134: Types of NAT Applications - PAT:

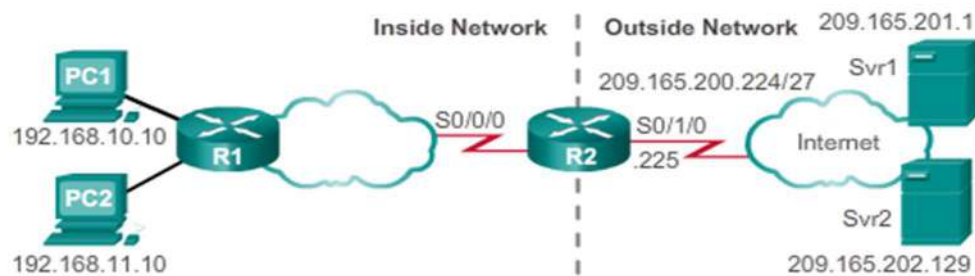
- PAT (also called NAT overload) allows the router to use one inside global address for many inside local addresses.
 - With address overloading, many privately addressed nodes can access the Internet using a single global address.
- There are two ways to configure PAT:
 - ISP allocates a single public IPv4 address
 - ISP allocates more than one public IPv4 address
- **Note:**
 - Over 65,000 inside addresses can theoretically map to a single outside address.
 - However, 4000 local addresses per global address is more realistic.
 - Each NAT translation consumes about 160 bytes of router DRAM.

NAT Overload (PAT)



- The NAT router keeps track of the different conversations by mapping TCP and UDP port numbers in the translation table.
 - Called an extended table entry.

Configuring PAT Using a Pool Example



```

R2(config)# ip nat pool NAT-POOL2 209.165.200.226 209.165.200.240
prefix-length 27
R2(config)# access-list 1 permit 192.168.0.0 0.0.255.255
R2(config)#
R2(config)# ip nat inside source list 1 pool NAT-POOL2 overload
R2(config)# interface Serial0/0/0
R2(config-if)# ip nat inside
R2(config-if)# exit
R2(config)# interface Serial0/1/0
R2(config-if)# ip nat outside

```

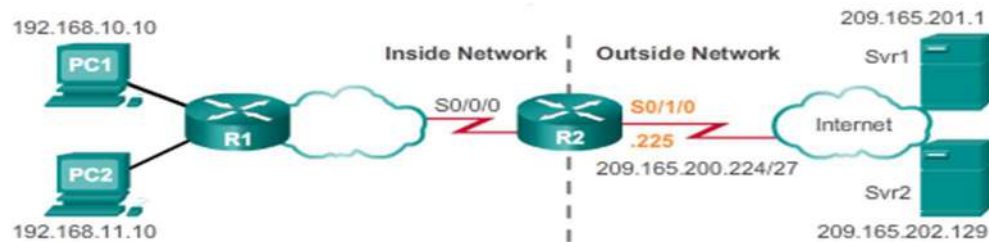
Verifying PAT Using a Pool Example

```

R2# show ip nat translations
Pro Inside global      Inside local      Outside local      Outside global
tcp 209.165.200.226:51839 192.168.10.10:51839 209.165.201.1:80 209.165.201.1:80
tcp 209.165.200.226:42558 192.168.11.10:42558 209.165.202.129:80 209.165.202.129:80
R2#

```

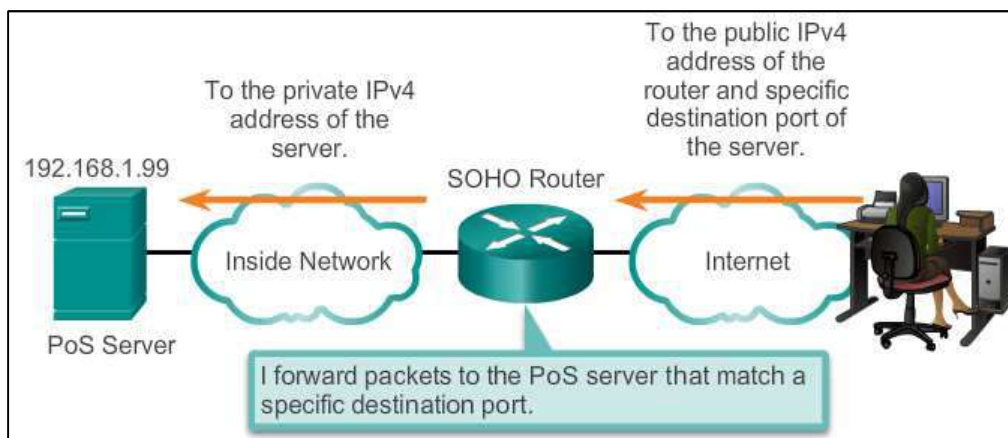
Configure PAT Using an Address Example



```
R2(config)# access-list 1 permit 192.168.0.0 0.0.255.255
R2(config)#
R2(config)# ip nat source list 1 interface serial 0/1/0 overload
R2(config)#
R2(config)# interface Serial0/0/0
R2(config-if)# ip nat inside
R2(config-if)# exit
R2(config)# interface Serial0/1/0
R2(config-if)# ip nat outside
```

Topic 135: Port Forwarding:

- Port forwarding (sometimes referred to as tunneling) is the act of forwarding traffic addressed to a specific network port from one network node to another.
 - Helpful in situations where servers have private addresses, not reachable from the outside networks.
 - Port forwarding can be enabled for applications by specifying the inside local address that requests should be forwarded to.



Configuring Port Forwarding with IOS

- In IOS, Port forwarding is essentially a static NAT translation with a specified TCP or UDP port number.

- Configured using the **ip nat inside source {static {tcp | udp local-ip local-port global-ip global-port} [extendable]** global configuration command.

Parameter	Description
tcp or udp	Indicates if this is a TCP or UDP port number.
<i>local-ip</i>	This is the IPv4 address assigned to the inside host (typically a private address).
<i>local-port</i>	- Sets the local TCP/UDP port in a range from 1-65535. - This is the port number the server is listening on.
<i>global-ip</i>	- Sets the global TCP/UDP port in a range from 1-65535. - This is the port number the outside client will use to reach the internal server.
extendable	The option is applied by default and allows the router to extend the translation to more than one port if necessary.

Topic 136: Packet Tracer – Configuring NAT Pool Overload and PAT-1:

- Build a Network
- Configure NAT Pool Overload
- Configure PAT

Topic 137: Packet Tracer – Configuring NAT Pool Overload and PAT-2:

This is the continuation of the previous topic.

Topic 138: Packet Tracer – Configuring NAT Pool Overload and PAT-3:

This is the continuation of the previous topic.

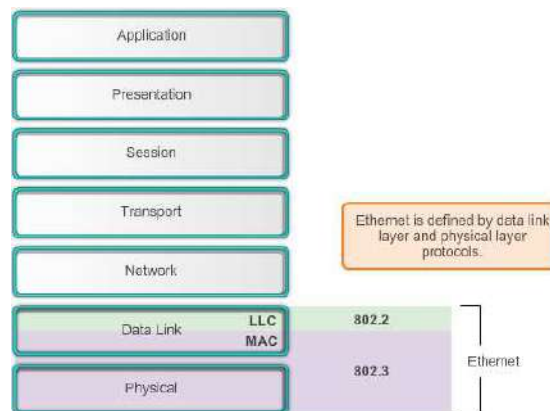
Topic 139: Ethernet Protocol:

Ethernet is the most widely used LAN technology used today. Ethernet operates in the data link layer and the physical layer. It is a family of networking technologies that are defined in the IEEE 802.2 and 802.3 standards. Ethernet supports data bandwidths of:

- 10 Mb/s
- 100 Mb/s
- 1000 Mb/s (1 Gb/s)

- 10,000 Mb/s (10 Gb/s)
- 40,000 Mb/s (40 Gb/s)
- 100,000 Mb/s (100 Gb/s)

As shown in Figure below, Ethernet standards define both the Layer 2 protocols and the Layer 1 technologies. For the Layer 2 protocols, as with all 802 IEEE standards, Ethernet relies on the two separate sublayers of the data link layer to operate, the Logical Link Control (LLC) and the MAC sublayers.



LLC sublayer

The Ethernet LLC sublayer handles the communication between the upper layers and the lower layers. This is typically between the networking software and the device hardware. The LLC sublayer takes the network protocol data, which is typically an IPv4 packet, and adds control information to help deliver the packet to the destination node. The LLC is used to communicate with the upper layers of the application, and transition the packet to the lower layers for delivery.

LLC is implemented in software, and its implementation is independent of the hardware. In a computer, the LLC can be considered the driver software for the NIC. The NIC driver is a program that interacts directly with the hardware on the NIC to pass the data between the MAC sublayer and the physical media.

MAC sublayer

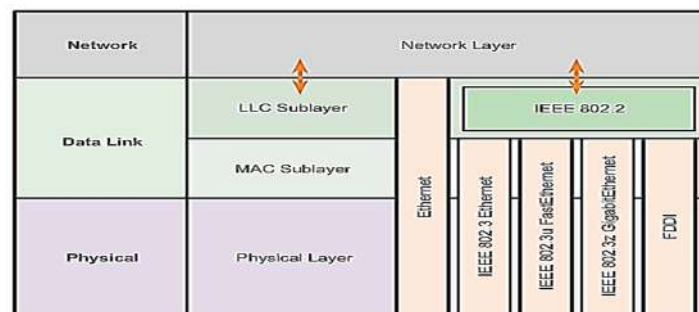
MAC constitutes the lower sublayer of the data link layer. MAC is implemented by hardware, typically in the computer NIC. The specifics are specified in the IEEE 802.3 standards. Figure below lists common IEEE Ethernet standards.

Ethernet MAC sublayer has two primary responsibilities:

- Data encapsulation
- Media access control

Data encapsulation

The data encapsulation process includes frame assembly before transmission, and frame disassembly upon reception of a frame. In forming the frame, the MAC layer adds a header and trailer to the network layer PDU.



Data encapsulation provides three primary functions:

Frame delimiting: The framing process provides important delimiters that are used to identify a group of bits that make up a frame. This process provides synchronization between the transmitting and receiving nodes.

Addressing: The encapsulation process also provides for data link layer addressing. Each Ethernet header added in the frame contains the physical address (MAC address) that enables a frame to be delivered to a destination node.

Error detection: Each Ethernet frame contains a trailer with a cyclic redundancy check (CRC) of the frame contents. After reception of a frame, the receiving node creates a CRC to compare to the one in the frame. If these two CRC calculations match, the frame can be trusted to have been received without error.

The use of frames aids in the transmission of bits as they are placed on the media and in the grouping of bits at the receiving node.

Media Access Control

The second responsibility of the MAC sublayer is media access control. Media access control is responsible for the placement of frames on the media and the removal of frames from the media. As its name implies, it controls access to the media. This sublayer communicates directly with the physical layer.

The underlying logical topology of Ethernet is a multi-access bus; therefore, all nodes (devices) on a single network segment share the medium. Ethernet is a contention-based method of networking. Recall that a contention-based method, or non-deterministic method, means that any device can try to transmit data across the shared medium whenever it has data to send. However, much like if two people try to talk simultaneously, if multiple devices on a single medium attempt to forward data simultaneously, the data will collide resulting in corrupted, unusable data. For this reason, Ethernet provides a method for controlling how the nodes share access through the use of a Carrier Sense Multiple Access (CSMA) technology.

Media Access Control

The CSMA process is used to first detect if the media is carrying a signal. If a carrier signal on the media from another node is detected, it means that another device is transmitting. When the device attempting to transmit sees that the media is busy, it will wait and try again after a short time period. If no carrier signal is detected, the device transmits its data. It is possible that the CSMA process will fail and two devices will transmit at the same time. This is called a data collision. If this occurs, the data sent by both devices will be corrupted and will need to be resent.

Contention-based media access control methods do not require mechanisms for tracking whose turn it is to access the media; therefore, they do not have the overhead of controlled access methods. However, the contention-based systems do not scale well under heavy media use. As use and the number of nodes increases, the probability of successful media access without a collision decreases. Additionally, the recovery mechanisms required to correct errors due to these collisions further diminishes the throughput.

CSMA is usually implemented in conjunction with a method for resolving media contention. The two commonly used methods are:

CSMA/Collision Detection

In CSMA/Collision Detection (CSMA/CD), the device monitors the media for the presence of a data signal. If a data signal is absent, indicating that the media is free, the device transmits the data. If signals are then detected that show another device was transmitting at the same time, all devices stop sending and try again later. Traditional forms of Ethernet were developed to use this method.

The widespread incorporation of switched technologies in modern networks has largely displaced the original need for CSMA/CD in local-area networks. Almost all wired connections between devices in a LAN today are full-duplex connections - a device is able to send and receive simultaneously. This means, that while Ethernet networks are designed with CSMA/CD technology, with today's intermediate devices, collisions do not occur and the processes utilized by CSMA/CD are really unnecessary.

However, wireless connections in a LAN environment still have to take collisions into account. Wireless LAN devices utilize the CSMA/Collision Avoidance (CSMA/CA) media access method.

CSMA/Collision Avoidance

In CSMA/CA, the device examines the media for the presence of a data signal. If the media is free, the device sends a notification across the media of its intent to use it. The device then sends the data. This method is used by 802.11 wireless networking technologies.

MAC Address: Ethernet Identity

The underlying logical topology of Ethernet is a multi-access bus. Every network device is connected to the same, shared media, and all the nodes receive all frames transmitted. The issue is if all devices are receiving every frame, how can each individual device identify if it is the intended receiver without the overhead of having to process and de-encapsulate the frame to get to the IP address? The issue becomes even more problematic in large, high traffic volume networks where lots of frames are forwarded.

To prevent the excessive overhead involved in the processing of every frame, a unique identifier called a MAC address was created to identify the actual source and destination nodes within an Ethernet network. Regardless of which variety of Ethernet is used, MAC addressing provided a method for device identification at the lower level of the OSI model. As you may recall, MAC addressing is added as part of a Layer 2 PDU. An Ethernet MAC address is a 48-bit binary value expressed as 12 hexadecimal digits (4 bits per hexadecimal digit).

MAC Address Structure

MAC addresses must be globally unique. The MAC address value is a direct result of IEEE-enforced rules for vendors to ensure globally unique addresses for each Ethernet device. The rules established by IEEE require any vendor that sells Ethernet devices to register with IEEE. The IEEE assigns the vendor a 3-byte (24-bit) code, called the Organizationally Unique Identifier (OUI).

IEEE requires a vendor to follow two simple rules:

- All MAC addresses assigned to a NIC or other Ethernet device must use that vendor's assigned OUI as the first 3 bytes.
- All MAC addresses with the same OUI must be assigned a unique value (vendor code or serial number) in the last 3 bytes.

Topic 140: Ethernet MAC:

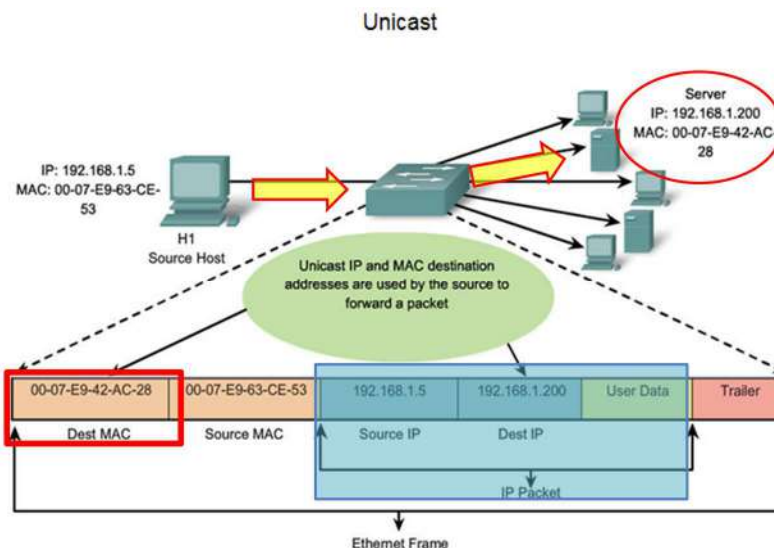
MAC Address

- Layer 2 Ethernet MAC address is a 48-bit binary value expressed as 12 hexadecimal digits

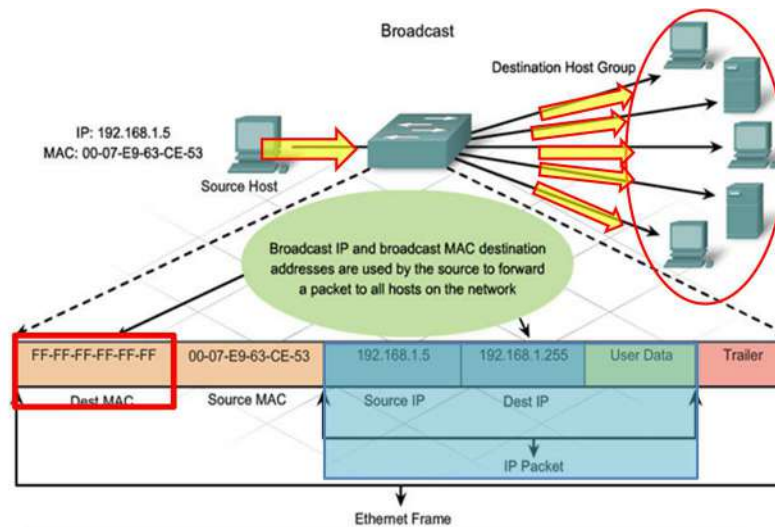
OUI unique

An Intel MAC address: **00-21-CC-BA-44-C4**

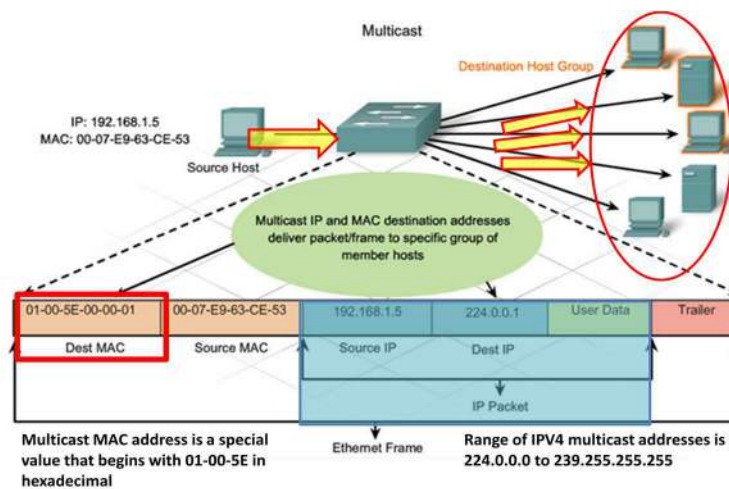
Unicast MAC Address



Broadcast MAC Address



Multicast MAC Address



MAC and IP

There are two primary addresses assigned to a host device:

- Physical address (the MAC address)
- Logical address (the IP address)

Both the MAC address and IP address work together to identify a device on the network. The process of using the MAC address and the IP address to find a computer is similar to the process of using a name and address of an individual to send a letter.

A person's name usually does not change. A person's address on the other hand, relates to where they live and can change.

Similar to the name of a person, the MAC address on a host does not change; it is physically assigned to the host NIC and is known as the physical address. The physical address remains the same regardless of where the host is placed.

The IP address is similar to the address of a person. This address is based on where the host is actually located. Using this address, it is possible for a frame to determine the location of where a frame should be sent. The IP address, or network address, is known as a logical address because it is assigned logically. It is assigned to each host by a network administrator based on the local network that the host is connected to. Both the physical MAC and logical IP addresses are required for a computer to communicate on a hierarchical network, just like both the name and address of a person are required to send a letter.

End-to-End Connectivity: MAC and IP

A source device will send a packet based on an IP address. One of the most common ways a source device determines the IP address of a destination device is through Domain Name Service (DNS), in which an IP address is associated to a domain name. For example, `www.cisco.com` is equal to `209.165.200.225`. This IP address will get the packet to the network location of the destination device. It is this IP address that routers will use to determine the best path to reach a destination. So, in short, IP addressing determines the end-to-end behavior of an IP packet.

However, along each link in a path, an IP packet is encapsulated in a frame specific to the particular data link technology associated with that link, such as Ethernet. End devices on an Ethernet network do not accept and process frames based on IP addresses, rather, a frame is accepted and processed based on MAC addresses.

On Ethernet networks, MAC addresses are used to identify, at a lower level, the source and destination hosts. When a host on an Ethernet network communicates, it sends frames containing its own MAC address as the source and the MAC address of the intended recipient as the destination. All hosts that receive the frame will read the destination MAC address. If the destination MAC address matches the MAC address configured on the host NIC, only then will the host process the message.

How are the IP addresses of the IP packets in a data flow associated with the MAC addresses on each link along the path to the destination? This is done through a process called Address Resolution Protocol (ARP).

Topic 141: Address Resolution Protocol (ARP):

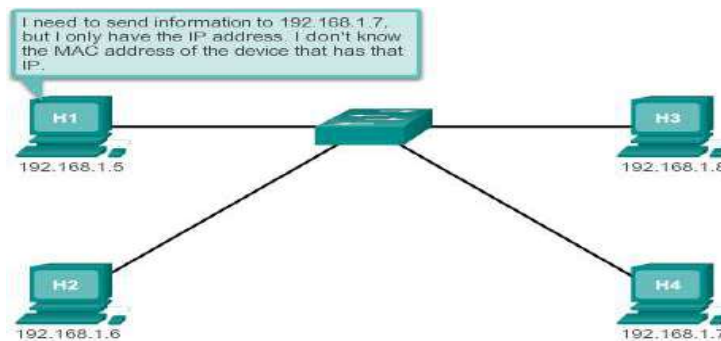
Recall that each node on an IP network has both a MAC address and an IP address. In order to send data, the node must use both of these addresses. The node must use its own MAC and IP addresses in the source fields and must provide both a MAC address and an IP address for the destination. While the IP address of the destination will be provided by a higher OSI layer, the sending node needs a way to find the MAC address of the destination for a given Ethernet link. This is the purpose of ARP.

ARP relies on certain types of Ethernet broadcast messages and Ethernet unicast messages, called ARP requests and ARP replies.

The ARP protocol provides two basic functions:

- Resolving IPv4 addresses to MAC addresses
- Maintaining a table of mappings

ARP Functions



ARP Table –

- Used to find data link layer address that is mapped to destination IPv4 address
- As a node receives frames from the media, it records the source IP and MAC address as a mapping in the ARP table

ARP request –

- Layer 2 broadcast to all devices on the Ethernet LAN
- The node that matches the IP address in the broadcast will reply
- If no device responds, the packet is dropped because a frame cannot be created

ARP Process – Communicating Locally

Creating the Frame

What does a node do when it needs to create a frame and the ARP cache does not contain a map of an IP address to a destination MAC address? It generates an ARP request!

When ARP receives a request to map an IPv4 address to a MAC address, it looks for the cached map in its ARP table. If an entry is not found, the encapsulation of the IPv4 packet fails and the Layer 2 processes notify ARP that it needs a map. The ARP processes then send out an ARP request packet to discover the MAC address of the destination device on the local network. If a device receiving the request has the destination IP address, it responds with an ARP reply. A map is created in the ARP table. Packets for that IPv4 address can now be encapsulated in frames.

If no device responds to the ARP request, the packet is dropped because a frame cannot be created. This encapsulation failure is reported to the upper layers of the device. If the device is an intermediary device, like a router, the upper layers may choose to respond to the source host with an error in an ICMPv4 packet.

ARP Process – Communicating Remotely

All frames must be delivered to a node on the local network segment. If the destination IPv4 host is on the local network, the frame will use the MAC address of this device as the destination MAC address.

If the destination IPv4 host is not on the local network, the source node needs to deliver the frame to the router interface that is the gateway or next hop used to reach that destination. The source node will use the MAC address of the gateway as the destination address for frames containing an IPv4 packet addressed to hosts on other networks.

The gateway address of the router interface is stored in the IPv4 configuration of the hosts. When a host creates a packet for a destination, it compares the destination IP address and its own IP address to determine if the two IP addresses are located on the same Layer 3 network. If the receiving host is not on the same network, the source uses the ARP process to determine a MAC address for the router interface serving as the gateway.

In the event that the gateway entry is not in the table, the normal ARP process will send an ARP request to retrieve the MAC address associated with the IP address of the router interface.

Removing Entries from ARP Table

For each device, an ARP cache timer removes ARP entries that have not been used for a specified period of time. The times differ depending on the device and its operating system. For example, some Windows operating systems store ARP cache entries for 2 minutes. If the entry is used again during that time, the ARP timer for that entry is extended to 10 minutes.

Commands may also be used to manually remove all or some of the entries in the ARP table. After an entry has been removed, the process for sending an ARP request and receiving an ARP reply must occur again to enter the map in the ARP table.

Each device has an operating system-specific command to delete the contents of the ARP cache. These commands do not invoke the execution of ARP in any way. They merely remove the entries of the ARP table. ARP service is integrated within the IPv4 protocol and implemented by the device. Its operation is transparent to both upper layer applications and users.

ARP Issues and Mitigations

Broadcast and security issues related to ARP can be mitigated with modern switches. Cisco switches support several security technologies specifically designed to mitigate Ethernet issues related to broadcasts, in general, and ARP, in particular.

Switches provide segmentation of a LAN, dividing the LAN into independent collision domains. Each port on a switch represents a separate collision domain and provides the full media bandwidth to the node or nodes connected on that port. While switches do not by default prevent broadcasts from propagating to connected devices, they do isolate unicast Ethernet communications so that they are only "heard" by the source and destination devices. So if there are a large number of ARP requests, each ARP reply will only be between two devices.

Topic 142: Switched Networks:

The role of switched networks has evolved dramatically in the last two decades. It was not long ago that flat Layer 2 switched networks were the norm. Flat Layer 2 data networks relied on the basic properties of Ethernet and the widespread use of hub repeaters to propagate LAN traffic throughout an

organization. Networks have fundamentally changed to switched LANs in a hierarchical network. A switched LAN allows more flexibility, traffic management, and additional features, such as:

- Quality of service
- Additional security
- Support for wireless networking and connectivity
- Support for new technologies, such as IP telephony and mobility services
- Layer 3 functionality

Switch Form Factor

There are various types of switches used in business networks. It is important to deploy the appropriate types of switches based on network requirements.

When selecting the type of switch, the network designer must choose between a fixed or a modular configuration, and stackable or non-stackable. Another consideration is the thickness of the switch, which is expressed in number of rack units. This is important for switches that are mounted in a rack.

Fixed Configuration Switches

Fixed configuration switches do not support features or options beyond those that originally came with the switch. The particular model determines the features and options available. For example, a 24-port gigabit fixed switch cannot support additional ports. There are typically different configuration choices that vary in how many and what types of ports are included with a fixed configuration switch.

Modular Configuration Switches

Modular configuration switches offer more flexibility in their configuration. Modular configuration switches typically come with different sized chassis that allow for the installation of different numbers of modular line cards. The line cards actually contain the ports. The line card fits into the switch chassis the way that expansion cards fit into a PC. The larger the chassis, the more modules it can support. There can be many different chassis sizes to choose from. A modular switch with a single 24-port line card could have an additional 24-port line card added to bring the total number of ports up to 48.

Stackable Configuration Switches

Stackable configuration switches can be interconnected using a special cable that provides high-bandwidth throughput between the switches. Cisco StackWise technology allows the interconnection of up to nine switches. Switches can be stacked one on top of the other with cables connecting the switches in a daisy chain fashion. The stacked switches effectively operate as a single larger switch. Stackable switches are desirable where fault tolerance and bandwidth availability are critical and a modular switch is too costly to implement. Using cross-connected connections, the network can recover quickly if a single switch fails. Stackable switches use a special port for interconnections. Many Cisco stackable switches also support StackPower technology, which enables power sharing among stack members.

Topic 143: Switch Features:

Port Density

The port density of a switch refers to the number of ports available on a single switch.

Fixed configuration switches typically support up to 48 ports on a single device. They have options for up to four additional ports for small form-factor pluggable (SFP) devices. High-port densities allow for better use of limited space and power. If there are two switches that each contain 24 ports, they would be able to support up to 46 devices, because at least one port per switch is lost with the connection of each switch to the rest of the network. In addition, two power outlets are required. Alternatively, if there is a single 48-port switch, 47 devices can be supported, with only one port used to connect the switch to the rest of the network, and only one power outlet needed to accommodate the single switch.

Modular switches can support very high-port densities through the addition of multiple switch port line cards. For example, some Catalyst 6500 switches can support in excess of 1,000 switch ports.

Large enterprise networks that support many thousands of network devices require high density, modular switches to make the best use of space and power. Without using a high-density modular switch, the network would need many fixed configuration switches to accommodate the number of devices that need network access. This approach can consume many power outlets and a lot of closet space.

The network designer must also consider the issue of uplink bottlenecks. For example, to achieve target performance, a series of fixed configuration switches may require many ports for bandwidth aggregation between switches. With a single modular switch, bandwidth aggregation is less of an issue, because the backplane of the chassis can provide the necessary bandwidth to accommodate the devices connected to the switch port line cards.

Forwarding Rates

Forwarding rates define the processing capabilities of a switch by rating how much data the switch can process per second. Switch product lines are classified by forwarding rates. Entry-level switches have lower forwarding rates than enterprise-level switches. Forwarding rates are important to consider when selecting a switch. If the switch forwarding rate is too low, it cannot accommodate full wire-speed communication across all of its switch ports. Wire speed is the data rate that each Ethernet port on the switch is capable of attaining. Data rates can be 100 Mb/s, 1 Gb/s, 10 Gb/s, or 100 Gb/s.

For example, a typical 48-port Gigabit Ethernet switch operating at full wire speed generates 48 Gb/s of traffic. If the switch only supports a forwarding rate of 32 Gb/s, it cannot run at full wire speed across all ports simultaneously. Fortunately, access layer switches typically do not need to operate at full wire speed, because they are physically limited by their uplinks to the distribution layer. This means that less expensive, lower performing switches can be used at the access layer, and more expensive, higher performing switches can be used at the distribution and core layers, where the forwarding rate has a greater impact on network performance.

Power over Ethernet (PoE)

PoE allows the switch to deliver power to a device over the existing Ethernet cabling. This feature can be used by IP phones and some wireless access points.

PoE allows more flexibility when installing wireless access points and IP phones, allowing them to be installed anywhere that there is an Ethernet cable. A network administrator should ensure that the PoE features are required, because switches that support PoE are expensive.

The relatively new Cisco Catalyst 2960-C and 3560-C Series compact switches support PoE pass-through. PoE pass-through allows a network administrator to power PoE devices connected to the switch, as well as the switch itself, by drawing power from certain upstream switches.

Topic 144: Frame Forwarding:

The concept of switching and forwarding frames is universal in networking and telecommunications. Various types of switches are used in LANs, WANs, and the public switched telephone network (PSTN). The fundamental concept of switching refers to a device making a decision based on two criteria:

- Ingress port
- Destination address

The decision on how a switch forwards traffic is made in relation to the flow of that traffic. The term ingress is used to describe where a frame enters the device on a port. The term egress is used to describe frames leaving the device from a particular port.

When a switch makes a decision, it is based on the ingress port and the destination address of the message.

A LAN switch maintains a table that it uses to determine how to forward traffic through the switch. The only intelligence of the LAN switch is its ability to use its table to forward traffic based on the ingress port and the destination address of a message. With a LAN switch, there is only one master switching table that describes a strict association between addresses and ports; therefore, a message with a given destination address always exits the same egress port, regardless of the ingress port it enters.

Cisco LAN switches forward Ethernet frames based on the destination MAC address of the frames.

Dynamically Populating a MAC Address Table

Switches use MAC addresses to direct network communications through the switch to the appropriate port toward the destination. A switch is made up of integrated circuits and the accompanying software that controls the data paths through the switch. For a switch to know which port to use to transmit a frame, it must first learn which devices exist on each port. As the switch learns the relationship of ports to devices, it builds a table called a MAC address, or content addressable memory (CAM) table. CAM is a special type of memory used in high-speed searching applications.

LAN switches determine how to handle incoming data frames by maintaining the MAC address table. A switch builds its MAC address table by recording the MAC address of each device connected to each of its ports. The switch uses the information in the MAC address table to send frames destined for a specific device out the port which has been assigned to that device.

A switch populates the MAC address table based on source MAC addresses. When a switch receives an incoming frame with a destination MAC address that is not found in the MAC address table, the switch forwards the frame out of all ports (flooding) except for the ingress port of the frame. When the

destination device responds, the switch adds the source MAC address of the frame and the port where the frame was received to the MAC address table. In networks with multiple interconnected switches, the MAC address table contains multiple MAC addresses for a single port connected to the other switches.

Switch Forwarding Methods

As networks grew and enterprises began to experience slower network performance, Ethernet bridges (an early version of a switch) were added to networks to limit the size of the collision domains. In the 1990s, advancements in integrated circuit technologies allowed for LAN switches to replace Ethernet bridges. These LAN switches were able to move the Layer 2 forwarding decisions from software to application-specific-integrated circuits (ASICs). ASICs reduce the packet-handling time within the device, and allow the device to handle an increased number of ports without degrading performance. This method of forwarding data frames at Layer 2 was referred to as store-and-forward switching. This term distinguished it from cut-through switching.

The store-and-forward method makes a forwarding decision on a frame after it has received the entire frame and checked the frame for errors using a mathematical error-checking mechanism known as a cyclic redundancy check (CRC).

By contrast, the cut-through method, begins the forwarding process after the destination MAC address of an incoming frame and the egress port has been determined.

Topic 145: Switching Domains:

Collision Domains

In hub-based Ethernet segments, network devices compete for the medium, because devices must take turns when transmitting. The network segments that share the same bandwidth between devices are known as collision domains, because when two or more devices within that segment try to communicate at the same time, collisions may occur.

It is possible, however, to use a switch device, operating the OSI data link layer, to divide a network into segments and reduce the number of devices that compete for bandwidth. When a switch is used, each port represents a new segment. Each new segment is a new collision domain. More bandwidth is available to the devices on the segment, and collisions in one collision domain do not interfere with the other segments. This is also known as micro segmentation.

Broadcast Domains

Although switches filter most frames based on MAC addresses, they do not filter broadcast frames. For other devices on the LAN to receive broadcast frames, switches must flood these frames out all ports except the one on which the broadcast was received. A collection of interconnected switches forms a single broadcast domain. Only a network layer device, such as a router, can divide a Layer 2 broadcast domain. Routers are used to segment both collision and broadcast domains.

When a device sends a Layer 2 broadcast, the destination MAC address in the frame is set to all binary ones. A frame with a destination MAC address of all binary ones is received by all devices in the broadcast domain.

The Layer 2 broadcast domain is referred to as the MAC broadcast domain. The MAC broadcast domain consists of all devices on the LAN that receive broadcast frames from a host.

When a switch receives a broadcast frame, it forwards the frame out each of its ports, except the ingress port where the broadcast frame was received. Each device connected to the switch receives a copy of the broadcast frame and processes it. Broadcasts are sometimes necessary for initially locating other devices and network services, but they also reduce network efficiency. Network bandwidth is used to propagate the broadcast traffic. Too many broadcasts and a heavy traffic load on a network can result in congestion: a slow-down in the network performance.

Topic 146: Basic Switch Configuration:

Bootup Process

After a Cisco switch is powered on, it goes through the following boot sequence:

1. First, the switch loads a power-on self-test (POST) program stored in ROM. POST checks the CPU subsystem. It tests the CPU, DRAM, and the portion of the flash device that makes up the flash file system.
2. Next, the switch loads the boot loader software. The boot loader is a small program stored in ROM and is run immediately after POST successfully completes.
3. The boot loader performs low-level CPU initialization. It initializes the CPU registers, which control where physical memory is mapped, the quantity of memory, and its speed.
4. The boot loader initializes the flash file system on the system board.
5. Finally, the boot loader locates and loads a default IOS operating system software image into memory and hands control of the switch over to the IOS.

The boot loader finds the Cisco IOS image and attempts to automatically boot by using information in the BOOT environment variable. If this variable is not set, the switch attempts to load and execute the first executable file it can by performing a recursive, depth-first search throughout the flash file system. In a depth-first search of the file system, the search begins at the first top-level directory. The search proceeds through the directory from the lowest level subdirectory, up the tree. If the search is unsuccessful, the next top-level directory is located and the bottom up search pattern is repeated. On Catalyst 2960 Series switches, the image file is normally contained in a directory that has the same name as the image file (excluding the .bin file extension).

The IOS operating system then initializes the interfaces using the Cisco IOS commands found in the configuration file, startup-config, which is stored in NVRAM.

The BOOT environment variable is set using the **boot system** global configuration mode command. Notice that the IOS is located in a distinct folder and the folder path is specified. Use the **show bootvar** command (**show boot** in older IOS versions) to see what the current IOS boot file is set to.

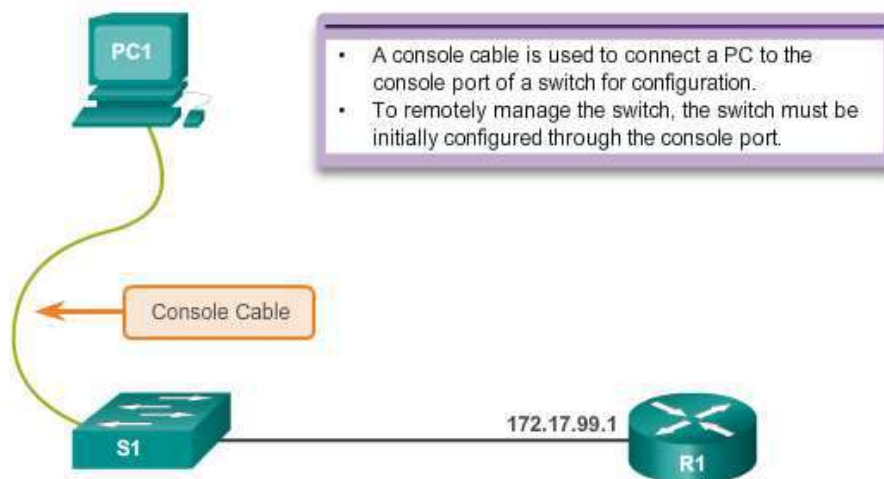
Preparing for Basic Switch Management

To prepare a switch for remote management access, the switch must be configured with an IP address and a subnet mask. Keep in mind, that to manage the switch from a remote network, the switch must be configured with a default gateway. This is very similar to configuring the IP address information on host devices. The SVI is a virtual interface, not a physical port on the switch.

SVI is a concept related to VLANs. VLANs are numbered logical groups to which physical ports can be assigned. Configurations and settings applied to a VLAN are also applied to all the ports assigned to that VLAN.

By default, the switch is configured to have the management of the switch controlled through VLAN 1. All ports are assigned to VLAN 1 by default. For security purposes, it is considered a best practice to use a VLAN other than VLAN 1 for the management VLAN.

Note that these IP settings are only for remote management access to the switch; the IP settings do not allow the switch to route Layer 3 packets.



Step 1. Configure Management Interface

An IP address and subnet mask is configured on the management SVI of the switch from VLAN interface configuration mode. As shown in Figure above, the **interface vlan 99** command is used to enter interface configuration mode. The **ip address** command is used to configure the IP address. The **no shutdown** command enables the interface. In this example, VLAN 99 is configured with IP address 172.17.99.11.

The SVI for VLAN 99 will not appear as "up/up" until VLAN 99 is created and there is a device connected to a switch port associated with VLAN 99. To create a VLAN with the vlan_id of 99, and associate it to an interface, use the following commands:

```
S1(config)# vlan vlan_id
```

```
S1(config-vlan)# name vlan_name
```

```
S1(config-vlan)# exit
```

```
S1(config)# interfaceinterface_id
```

```
S1(config-if)# switchport access vlan vlan_id
```

Step 2. Configure Default Gateway

The switch should be configured with a default gateway if it will be managed remotely from networks not directly connected. The default gateway is the router the switch is connected to. The switch will forward its IP packets with destination IP addresses outside the local network to the default gateway. As shown in Figure above, R1 is the default gateway for S1. The interface on R1 connected to the switch has IP address 172.17.99.1. This address is the default gateway address for S1.

To configure the default gateway for the switch, use the **ip default-gateway** command. Enter the IP address of the default gateway. The default gateway is the IP address of the router interface to which the switch is connected. Use the **copy running-config startup-config** command to back up your configuration.

Step 3. Verify Configuration

The **show ip interface brief** command is useful when determining the status of both physical and virtual interfaces. The output shown confirms that interface VLAN 99 has been configured with an IP address and subnet mask and that it is operational.

Topic 147: Configure Switch Ports:

Duplex Communication

Full-duplex communication improves the performance of a switched LAN. Full-duplex communication increases effective bandwidth by allowing both ends of a connection to transmit and receive data simultaneously. This is also known as bidirectional. This method of optimizing network performance requires micro-segmentation. A micro-segmented LAN is created when a switch port has only one device connected and is operating at full-duplex. This results in a micro size collision domain of a single device. However, because there is only one device connected, a micro-segmented LAN is collision free.

Unlike full-duplex communication, half-duplex communication is unidirectional. Sending and receiving data does not occur at the same time. Half-duplex communication creates performance issues because data can flow in only one direction at a time, often resulting in collisions. Half-duplex connections are typically seen in older hardware, such as hubs. Full-duplex communication has replaced half-duplex in most hardware.

Most Ethernet and Fast Ethernet NICs sold today offer full-duplex capability. Gigabit Ethernet and 10Gb NICs require full-duplex connections to operate. In full-duplex mode, the collision detection circuit on the NIC is disabled. Frames that are sent by the two connected devices cannot collide because the devices use two separate circuits in the network cable. Full-duplex connections require a switch that supports full-duplex configuration, or a direct connection using an Ethernet cable between two devices.

Duplex and Speed

Switch ports can be manually configured with specific duplex and speed settings. Use the **duplex** interface configuration mode command to manually specify the duplex mode for a switch

port. Use the **speed** interface configuration mode command to manually specify the speed for a switch port.

The default setting for both duplex and speed for switch ports on Cisco Catalyst 2960 and 3560 switches is auto. The 10/100/1000 ports operate in either half- or full-duplex mode when they are set to 10 or 100 Mb/s, but when they are set to 1000 Mb/s (1 Gb/s), they operate only in full-duplex mode. Auto-negotiation is useful when the speed and duplex settings of the device connecting to the port are unknown or may change. When connecting to known devices, such as servers, dedicated workstations, or network devices, best practice is to manually set the speed and duplex settings.

When troubleshooting switch port issues, the duplex and speed settings should be checked.

All fiber optic ports, such as 100BASE-FX ports, operate only at one preset speed and are always full-duplex.

Auto-MDIX

Until recently, certain cable types (straight-through or crossover) were required when connecting devices. Switch-to-switch or switch-to-router connections required using different Ethernet cables. Using the automatic medium-dependent interface crossover (auto-MDIX) feature on an interface eliminates this problem. When auto-MDIX is enabled, the interface automatically detects the required cable connection type (straight-through or crossover) and configures the connection appropriately. When connecting to switches without the auto-MDIX feature, straight-through cables must be used to connect to devices such as servers, workstations, or routers and crossover cables must be used to connect to other switches or repeaters.

With auto-MDIX enabled, either type of cable can be used to connect to other devices, and the interface automatically adjusts to communicate successfully. On newer Cisco routers and switches, the **mdix auto** interface configuration mode command enables the feature. When using auto-MDIX on an interface, the interface speed and duplex must be set to **auto** so that the feature operates correctly.

To examine the auto-MDIX setting for a specific interface, use the **show controllers ethernet-controller** command with the **phy** keyword. To limit the output to lines referencing auto-MDIX, use the **include Auto-MDIX** filter.

Verifying Switch Port Configuration

The **show interfaces** command is another commonly used command, which displays status and statistics information on the network interfaces of the switch. The **show interfaces** command is frequently used when configuring and monitoring network devices.

Topic 148: Packet Tracer - Configuring Basic Switch Settings:

- Configure Switch Settings
- Verify Switch Settings

Topic 149: Secure Remote Access - SSH:

Secure Shell (SSH) is a protocol that provides a secure (encrypted) management connection to a remote device. SSH should replace Telnet for management connections. Telnet is an older protocol that uses unsecure plaintext transmission of both the login authentication (username and password) and the data transmitted between the communicating devices. SSH provides security for remote connections by providing strong encryption when a device is authenticated (username and password) and also for the transmitted data between the communicating devices. SSH is assigned to TCP port 22. Telnet is assigned to TCP port 23.

To enable SSH on a Catalyst 2960 switch, the switch must be using a version of the IOS software including cryptographic (encrypted) features and capabilities. Use the **show version** command on the switch to see which IOS the switch is currently running.

Configuring SSH

Before configuring SSH, the switch must be minimally configured with a unique hostname and the correct network connectivity settings.

Step 1. Verify SSH support.

Use the **show ip ssh** command to verify that the switch supports SSH. If the switch is not running an IOS that supports cryptographic features, this command is unrecognized.

Step 2. Configure the IP domain.

Configure the IP domain name of the network using the **ip domain-name domain-name** global configuration mode command.

Step 3. Generate RSA key pairs.

Generating an RSA key pair automatically enables SSH. Use the **crypto key generate rsa** global configuration mode command to enable the SSH server on the switch and generate an RSA key pair. When generating RSA keys, the administrator is prompted to enter a modulus length. Cisco recommends a minimum modulus size of 1,024 bits. A longer modulus length is more secure, but it takes longer to generate and to use.

Step 4. Configure user authentication.

The SSH server can authenticate users locally or using an authentication server. To use the local authentication method, create a username and password pair using the **username username password password** global configuration mode command.

Step 5. Configure the vty lines.

Enable the SSH protocol on the vty lines using the **transport input ssh** line configuration mode command. The Catalyst 2960 has vty lines ranging from 0 to 15. This configuration prevents non-SSH (such as Telnet) connections and limits the switch to accept only SSH connections. Use the **line vty** global configuration mode command and then the **login local** line configuration mode command to require local authentication for SSH connections from the local username database.

Step 6. Enable SSH version 2.

By default, SSH supports both versions 1 and 2. When supporting both versions, this is shown in the **show ip ssh** output as supporting version 1.99. Version 1 has known vulnerabilities. For this reason, it is recommended to enable only version 2. Enable SSH version using the **ip ssh version 2** global configuration command.

Topic 150: VLANs:

Within a switched internetwork, VLANs provide segmentation and organizational flexibility. VLANs provide a way to group devices within a LAN. A group of devices within a VLAN communicate as if they were attached to the same wire. VLANs are based on logical connections, instead of physical connections.

VLANs allow an administrator to segment networks based on factors such as function, project team, or application, without regard for the physical location of the user or device. Devices within a VLAN act as if they are in their own independent network, even if they share a common infrastructure with other VLANs. Any switch port can belong to a VLAN, and unicast, broadcast, and multicast packets are forwarded and flooded only to end stations within the VLAN where the packets are sourced. Each VLAN is considered a separate logical network, and packets destined for stations that do not belong to the VLAN must be forwarded through a device that supports routing.

A VLAN creates a logical broadcast domain that can span multiple physical LAN segments. VLANs improve network performance by separating large broadcast domains into smaller ones. If a device in one VLAN sends a broadcast Ethernet frame, all devices in the VLAN receive the frame, but devices in other VLANs do not.

VLANs enable the implementation of access and security policies according to specific groupings of users. Each switch port can be assigned to only one VLAN (with the exception of a port connected to an IP phone or to another switch).

VLAN Benefits

User productivity and network adaptability are important for business growth and success. VLANs make it easier to design a network to support the goals of an organization. The primary benefits of using VLANs are as follows:

Security - Groups that have sensitive data are separated from the rest of the network, decreasing the chances of confidential information breaches.

Cost reduction - Cost savings result from reduced need for expensive network upgrades and more efficient use of existing bandwidth and uplinks.

Better performance - Dividing flat Layer 2 networks into multiple logical workgroups (broadcast domains) reduces unnecessary traffic on the network and boosts performance.

Shrink broadcast domains - Dividing a network into VLANs reduces the number of devices in the broadcast domain.

Improved IT staff efficiency - VLANs make it easier to manage the network because users with similar network requirements share the same VLAN. When a new switch is provisioned, all the policies and

procedures already configured for the particular VLAN are implemented when the ports are assigned. It is also easy for the IT staff to identify the function of a VLAN by giving it an appropriate name.

Simpler project and application management - VLANs aggregate users and network devices to support business or geographic requirements. Having separate functions makes managing a project or working with a specialized application easier; an example of such an application is an e-learning development platform for faculty.

Each VLAN in a switched network corresponds to an IP network; therefore, VLAN design must take into consideration the implementation of a hierarchical network-addressing scheme. Hierarchical network addressing means that IP network numbers are applied to network segments or VLANs in an orderly fashion that takes the network as a whole into consideration. Blocks of contiguous network addresses are reserved for and configured on devices in a specific area of the network.

VLAN Types

There are a number of distinct types of VLANs used in modern networks. Some VLAN types are defined by traffic classes. Other types of VLANs are defined by the specific function that they serve.

Data VLAN

A data VLAN is a VLAN that is configured to carry user-generated traffic. A VLAN carrying voice or management traffic would not be a data VLAN. It is common practice to separate voice and management traffic from data traffic. A data VLAN is sometimes referred to as a user VLAN. Data VLANs are used to separate the network into groups of users or devices.

Default VLAN

All switch ports become a part of the default VLAN after the initial boot up of a switch loading the default configuration. Switch ports that participate in the default VLAN are part of the same broadcast domain. This allows any device connected to any switch port to communicate with other devices on other switch ports. The default VLAN for Cisco switches is VLAN 1. VLAN 1 has all the features of any VLAN, except it cannot be renamed or deleted. By default, all Layer 2 control traffic is associated with VLAN 1.

Native VLAN

A native VLAN is assigned to an 802.1Q trunk port. Trunk ports are the links between switches that support the transmission of traffic associated with more than one VLAN. An 802.1Q trunk port supports traffic coming from many VLANs (tagged traffic), as well as traffic that does not come from a VLAN (untagged traffic). Tagged traffic refers to traffic that has a 4-byte tag inserted within the original Ethernet frame header, specifying the VLAN to which the frame belongs. The 802.1Q trunk port places untagged traffic on the native VLAN, which by default is VLAN 1.

Native VLANs are defined in the IEEE 802.1Q specification to maintain backward compatibility with untagged traffic common to legacy LAN scenarios. A native VLAN serves as a common identifier on opposite ends of a trunk link.

It is a best practice to configure the native VLAN as an unused VLAN, distinct from VLAN 1 and other VLANs. In fact, it is not unusual to dedicate a fixed VLAN to serve the role of the native VLAN for all trunk ports in the switched domain.

Management VLAN

A management VLAN is any VLAN configured to access the management capabilities of a switch. VLAN 1 is the management VLAN by default. To create the management VLAN, the switch virtual interface (SVI) of that VLAN is assigned an IP address and subnet mask, allowing the switch to be managed via HTTP, Telnet, SSH, or SNMP. Because the out-of-the-box configuration of a Cisco switch has VLAN 1 as the default VLAN, VLAN 1 would be a bad choice for the management VLAN.

In the past, the management VLAN for a 2960 switch was the only active SVI. On 15.x versions of the Cisco IOS for Catalyst 2960 Series switches, it is possible to have more than one active SVI. With Cisco IOS 15.x, the particular active SVI assigned for remote management must be documented. While theoretically a switch can have more than one management VLAN, having more than one increases exposure to network attacks.

Topic 151: VLANs in a Multi-Switched Environment:

VLAN Trunks

A trunk is a point-to-point link between two network devices that carries more than one VLAN. A VLAN trunk extends VLANs across an entire network. Cisco supports IEEE 802.1Q for coordinating trunks on Fast Ethernet, Gigabit Ethernet, and 10-Gigabit Ethernet interfaces.

VLANs would not be very useful without VLAN trunks. VLAN trunks allow all VLAN traffic to propagate between switches, so that devices which are in the same VLAN, but connected to different switches, can communicate without the intervention of a router.

A VLAN trunk does not belong to a specific VLAN; rather, it is a conduit for multiple VLANs between switches and routers. A trunk could also be used between a network device and server or other device that is equipped with an appropriate 802.1Q-capable NIC. By default, on a Cisco Catalyst switch, all VLANs are supported on a trunk port.

Network without VLANs

In normal operation, when a switch receives a broadcast frame on one of its ports, it forwards the frame out all other ports except the port where the broadcast was received.

Network with VLANs

When VLANs are implemented on a switch, the transmission of unicast, multicast, and broadcast traffic from a host in a particular VLAN are restricted to the devices that are in that VLAN.

Tagging Ethernet Frames for VLAN Identification

Catalyst 2960 Series switches are Layer 2 devices. They use the Ethernet frame header information to forward packets. They do not have routing tables. The standard Ethernet frame header does not contain information about the VLAN to which the frame belongs; thus, when Ethernet frames are placed on a

trunk, information about the VLANs to which they belong must be added. This process, called tagging, is accomplished by using the IEEE 802.1Q header, specified in the IEEE 802.1Q standard. The 802.1Q header includes a 4-byte tag inserted within the original Ethernet frame header, specifying the VLAN to which the frame belongs.

When the switch receives a frame on a port configured in access mode and assigned a VLAN, the switch inserts a VLAN tag in the frame header, recalculates the FCS, and sends the tagged frame out of a trunk port.

VLAN Tag Field Details

The VLAN tag field consists of a Type field, a Priority field, a Canonical Format Identifier field, and VLAN ID field:

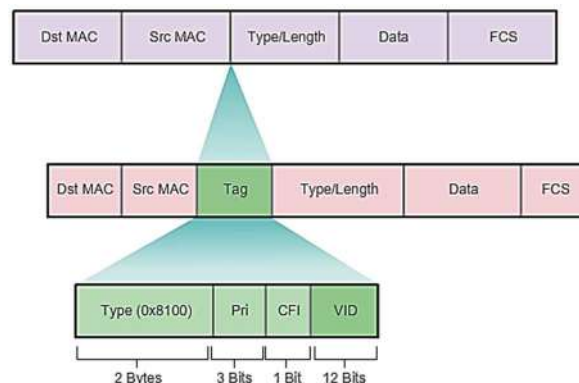
Type - A 2-byte value called the tag protocol ID (TPID) value. For Ethernet, it is set to hexadecimal 0x8100.

User priority - A 3-bit value that supports level or service implementation.

Canonical Format Identifier (CFI) - A 1-bit identifier that enables Token Ring frames to be carried across Ethernet links.

VLAN ID (VID) - A 12-bit VLAN identification number that supports up to 4096 VLAN IDs.

After the switch inserts the Type and tag control information fields, it recalculates the FCS values and inserts the new FCS into the frame.



Native VLANs

Tagged Frames on the Native VLAN

Some devices that support trunking, add a VLAN tag to native VLAN traffic. Control traffic sent on the native VLAN should not be tagged. If an 802.1Q trunk port receives a tagged frame with the VLAN ID the same as the native VLAN, it drops the frame. Consequently, when configuring a switch port on a Cisco switch, configure devices so that they do not send tagged frames on the native VLAN. Devices from other vendors that support tagged frames on the native VLAN include IP phones, servers, routers, and non-Cisco switches.

Untagged Frames on the Native VLAN

When a Cisco switch trunk port receives untagged frames (which are unusual in a well-designed network), it forwards those frames to the native VLAN. If there are no devices associated with the native VLAN (which is not unusual) and there are no other trunk ports (which is not unusual), then the frame is dropped. The default native VLAN is VLAN 1. When configuring an 802.1Q trunk port, a default Port VLAN ID (PVID) is assigned the value of the native VLAN ID. All untagged traffic coming in or out of the 802.1Q port is forwarded based on the PVID value. For example, if VLAN 99 is configured as the native VLAN, the PVID is 99 and all untagged traffic is forwarded to VLAN 99. If the native VLAN has not been reconfigured, the PVID value is set to VLAN 1.

Topic 152: VLAN Implementation:

Normal Range VLANs

Different Cisco Catalyst switches support various numbers of VLANs. The number of supported VLANs is large enough to accommodate the needs of most organizations. For example, the Catalyst 2960 and 3560 Series switches support over 4,000 VLANs. Normal range VLANs on these switches are numbered 1 to 1,005 and extended range VLANs are numbered 1,006 to 4,094.

Normal Range VLANs

- Used in small- and medium-sized business and enterprise networks.
- Identified by a VLAN ID between 1 and 1005.
- IDs 1002 through 1005 are reserved for Token Ring and FDDI VLANs.
- IDs 1 and 1002 to 1005 are automatically created and cannot be removed.
- Configurations are stored within a VLAN database file, called vlan.dat. The vlan.dat file is located in the flash memory of the switch.
- The VLAN Trunking Protocol (VTP), which helps manage VLAN configurations between switches, can only learn and store normal range VLANs.

Extended Range VLANs

- Enable service providers to extend their infrastructure to a greater number of customers. Some global enterprises could be large enough to need extended range VLAN IDs.
- Are identified by a VLAN ID between 1006 and 4094.
- Configurations are not written to the vlan.dat file.
- Support fewer VLAN features than normal range VLANs.
- Are, by default, saved in the running configuration file.
- VTP does not learn extended range VLANs.

Creating a VLAN

When configuring normal range VLANs, the configuration details are stored in flash memory on the switch in a file called vlan.dat. Flash memory is persistent and does not require the **copy running-config startup-config** command. However, because other details are often configured on a Cisco switch at the same time that VLANs are created, it is good practice to save running configuration changes to the startup configuration.

In addition to entering a single VLAN ID, a series of VLAN IDs can be entered separated by commas, or a range of VLAN IDs separated by hyphens using the **vlan** *vlan-id* command. For example, use the following command to create VLANs 100, 102, 105, 106, and 107:

```
S1(config)# vlan 100,102,105-107
```

Assigning Ports to VLANs

After creating a VLAN, the next step is to assign ports to the VLAN. An access port can belong to only one VLAN at a time; one exception to this rule is that of a port connected to an IP phone, in which case, there are two VLANs associated with the port: one for voice and one for data.

The **switchport access vlan** command forces the creation of a VLAN if it does not already exist on the switch.

Deleting a VLAN

The **no vlan** *vlan-id* global configuration mode command is used to remove VLAN 20 from the switch.

Caution: Before deleting a VLAN, be sure to first reassign all member ports to a different VLAN. Any ports that are not moved to an active VLAN are unable to communicate with other hosts after the VLAN is deleted and until they are assigned to an active VLAN.

Alternatively, the entire `vlan.dat` file can be deleted using the **delete flash: vlan.dat** privileged EXEC mode command. The abbreviated command version (**delete vlan.dat**) can be used if the `vlan.dat` file has not been moved from its default location. After issuing this command and reloading the switch, the previously configured VLANs are no longer present. This effectively places the switch into its factory default condition concerning VLAN configurations.

Note: For a Catalyst switch, the **erase startup-config** command must accompany the **delete vlan.dat** command prior to reload to restore the switch to its factory default condition.

Verifying VLAN Information

After a VLAN is configured, VLAN configurations can be validated using Cisco IOS **show** commands.

The **show interfaces vlan** *vlan-id* command displays details that are beyond the scope of this course.

Topic 153: VLAN Trunks:

Native VLANs

A VLAN trunk is an OSI Layer 2 link between two switches that carries traffic for all VLANs (unless the allowed VLAN list is restricted manually or dynamically). To enable trunk links, configure the ports on either end of the physical link with parallel sets of commands.

To configure a switch port on one end of a trunk link, use the **switchport mode trunk** command. With this command, the interface changes to permanent trunking mode. The port enters into a Dynamic Trunking Protocol (DTP) negotiation to convert the link into a trunk link even if the interface connecting to it does not agree to the change. DTP is described in the next topic. In this course, the **switchport mode trunk** command is the only method implemented for trunk configuration.

Use the Cisco IOS **switchport trunk allowed vlan** *vlan-list* command to specify the list of VLANs to be allowed on the trunk link.

Configuring IEEE 802.1Q Trunk Links

Cisco Switch IOS Commands	
Enter global configuration mode.	S1# configure terminal
Enter interface configuration mode.	S1(config)# interface interface_id
Force the link to be a trunk link.	S1(config-if)# switchport mode trunk
Specify a native VLAN for untagged 802.1Q trunks.	S1(config-if)# switchport trunk native vlan vlan_id
Specify the list of VLANs to be allowed on the trunk link.	S1(config-if)# switchport trunk allowed vlan vlan-list
Return to the privileged EXEC mode.	S1(config-if)# end

Resetting the Trunk to Default State

Cisco Switch IOS Commands	
Enter global configuration mode.	S1# configure terminal
Enter interface configuration mode.	S1(config)# interface interface_id
Set trunk to allow all VLANs.	S1(config-if)# no switchport trunk allowed vlan
Reset native VLAN to default.	S1(config-if)# no switchport trunk native vlan
Return to the privileged EXEC mode.	S1(config-if)# end

Topic 154: Dynamic Trunking Protocol (DTP):

Ethernet trunk interfaces support different trunking modes. An interface can be set to trunking or nontrunking, or to negotiate trunking with the neighbor interface. Trunk negotiation is managed by the Dynamic Trunking Protocol (DTP), which operates on a point-to-point basis only, between network devices.

DTP is a Cisco proprietary protocol that is automatically enabled on Catalyst 2960 and Catalyst 3560 Series switches. Switches from other vendors do not support DTP. DTP manages trunk negotiation only if the port on the neighbor switch is configured in a trunk mode that supports DTP.

To enable trunking from a Cisco switch to a device that does not support DTP, use the **switchport mode trunk** and **switchport nonegotiate** interface configuration mode commands. This causes the interface to become a trunk, but not generate DTP frames.

To Clear A Switch

S1# **delete vlan.dat**

Delete filename [vlan.dat]?

Delete flash:/vlan.dat? [confirm]

S1# **erase startup-config**

Erasing the nvram filesystem will remove all configuration files! Continue? [confirm]

[OK]

Erase of nvram: complete

%SYS-7-NV_BLOCK_INIT: Initialized the geometry of nvram

S1# **reload**

Proceed with reload? [confirm]

Topic 155: Packet Tracer – Configuring VLANs and Trunking-1:

- Build the Network
- Create VLANs and Assign Switch Ports
- Configure an 802.1Q Trunk between Switches

Topic 156: Packet Tracer – Configuring VLANs and Trunking-2:

This is the continuation of the previous topic.

Topic 157: Purpose of Spanning Tree:

Redundancy in a Hierarchical Network

- Network redundancy is a key to maintaining network reliability. Multiple physical links between devices provide redundant paths. The network can then continue to operate when a single link or port has failed. Redundant links can also share the traffic load and increase capacity.
- Multiple paths must be managed so that Layer 2 loops are not created, the best paths are chosen for use, and an alternate path is immediately available should a primary path fail. The Spanning Tree Protocols are used to manage Layer 2 redundancy.
- Redundant devices, such as multilayer switches or routers, provide the capability for a client to use an alternate default gateway should the primary default gateway fail. A client may now have multiple paths to more than one possible default gateway. First Hop Redundancy Protocols are used to manage how a client is assigned a default gateway, and to be able to use an alternate default gateway should the primary default gateway fail.

Issues with Layer 1 Redundancy–MAC DB Instability

MAC Database Instability

Ethernet frames do not have a time to live (TTL) attribute, like IP packets. As a result, if there is no mechanism enabled to block continued propagation of these frames on a switched network, they continue to propagate between switches endlessly, or until a link is disrupted and breaks the loop. This continued propagation between switches can result in MAC database instability. This can occur due to broadcast frames forwarding.

Broadcast frames are forwarded out all switch ports, except the original ingress port. This ensures that all devices in a broadcast domain are able to receive the frame. If there is more than one path for the frame to be forwarded out, an endless loop can result. When a loop occurs, it is possible for the MAC address table on a switch to constantly change with the updates from the broadcast frames, resulting in MAC database instability.

Broadcast Storm

A broadcast storm occurs when there are so many broadcast frames caught in a Layer 2 loop that all available bandwidth is consumed. Consequently, no bandwidth is available for legitimate traffic and the network becomes unavailable for data communication. This is an effective denial of service.

A broadcast storm is inevitable on a looped network. As more devices send broadcasts over the network, more traffic is caught within the loop, consuming resources. This eventually creates a broadcast storm that causes the network to fail.

There are other consequences of broadcast storms. Because broadcast traffic is forwarded out every port on a switch, all connected devices have to process all broadcast traffic that is being flooded endlessly around the looped network. This can cause the end device to malfunction because of the high processing requirements for sustaining such a high traffic load on the NIC.

Because devices connected to a network are regularly sending out broadcast frames, such as ARP requests, a broadcast storm can develop in seconds. As a result, when a loop is created, the switched network is quickly brought down.

Multiple Frame Transmissions

Broadcast frames are not the only type of frames that are affected by loops. Unicast frames sent onto a looped network can result in duplicate frames arriving at the destination device.

Most upper layer protocols are not designed to recognize, or cope with, duplicate transmissions. In general, protocols that make use of a sequence-numbering mechanism assume that the transmission has failed and that the sequence number has recycled for another communication session. Other protocols attempt to hand the duplicate transmission to the appropriate upper layer protocol to be processed and possibly discarded.

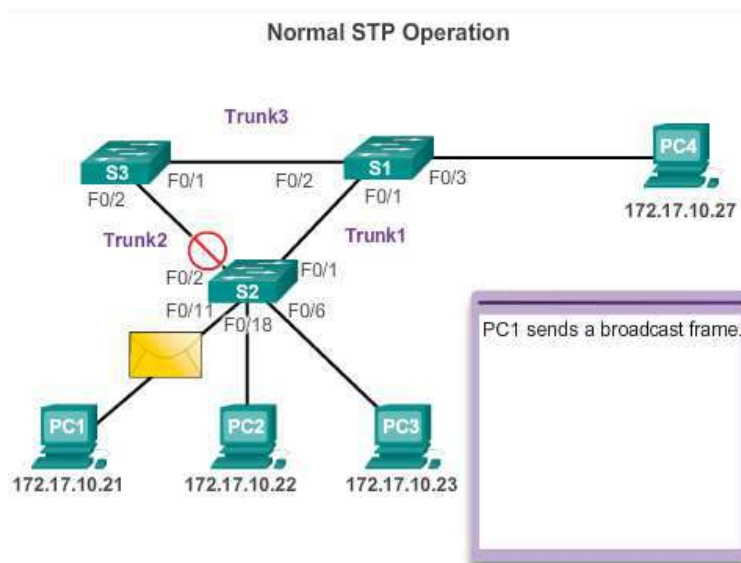
Layer 2 LAN protocols, such as Ethernet, lack a mechanism to recognize and eliminate endlessly looping frames. Some Layer 3 protocols implement a TTL mechanism that limits the number of times a Layer 3 networking device can retransmit a packet. Lacking such a mechanism, Layer 2 devices continue to retransmit looping traffic indefinitely. A Layer 2 loop-avoidance mechanism, STP, was developed to address these problems.

To prevent these issues from occurring in a redundant network, some type of spanning tree must be enabled on the switches. Spanning tree is enabled, by default, on Cisco switches to prevent Layer 2 loops from occurring.

Topic 158: Spanning Tree Protocol (STP) – Part I:

Redundancy increases the availability of the network topology by protecting the network from a single point of failure, such as a failed network cable or switch. When physical redundancy is introduced into a design, loops and duplicate frames occur. Loops and duplicate frames have severe consequences for a switched network. The STP was developed to address these issues.

STP ensures that there is only one logical path between all destinations on the network by intentionally blocking redundant paths that could cause a loop. A port is considered blocked when user data is prevented from entering or leaving that port. This does not include bridge protocol data unit (BPDU) frames that are used by STP to prevent loops. Blocking the redundant paths is critical to preventing loops on the network. The physical paths still exist to provide redundancy, but these paths are disabled to prevent the loops from occurring. If the path is ever needed to compensate for a network cable or switch failure, STP recalculates the paths and unblocks the necessary ports to allow the redundant path to become active.



In the example above, all switches have STP enabled:

1. PC1 sends a broadcast out onto the network.
2. S2 is configured with STP and has set the port for Trunk2 to a blocking state. The blocking state prevents ports from being used to forward user data, thus preventing a loop from occurring. S2 forwards a broadcast frame out all switch ports, except the originating port from PC1 and the port for Trunk2.
3. S1 receives the broadcast frame and forwards it out all of its switch ports, where it reaches PC4 and S3. S3 forwards the frame out the port for Trunk2 and S2 drops the frame. The Layer 2 loop is prevented.

STP prevents loops from occurring by configuring a loop-free path through the network using strategically placed "blocking-state" ports. The switches running STP are able to compensate for failures by dynamically unblocking the previously blocked ports and permitting traffic to traverse the alternate paths.

Up to now, we have used the term Spanning Tree Protocol and the acronym STP. The usage of the Spanning Tree Protocol term and the STP acronym can be misleading. Technically, these refer to the original 802.1D standard and its implementation. However, many professionals generically use these to refer to various implementations of spanning tree, such as Rapid Spanning Tree Protocol (RSTP) and Multiple Spanning Tree Protocol (MSTP). In order to communicate spanning tree concepts correctly, it is important to refer to the particular implementation or standard in context. The latest IEEE documentation on spanning tree, IEEE-802-1D-2004, says "STP has now been superseded by the Rapid Spanning Tree Protocol (RSTP)"; so one sees that the IEEE uses "STP" to refer to the original implementation of spanning tree and "RSTP" to describe the version of spanning tree specified in IEEE-802.1D-2004. In this curriculum, when the original Spanning Tree Protocol is the context of a discussion, the phrase original 802.1D spanning tree or simply 802.1D is used to avoid confusion.

Topic 159: Spanning Tree Protocol (STP) – Part II:

Spanning Tree Algorithm – Port Rules

The Spanning Tree Algorithm (STA) is used to determine which switch ports on a network must be put in blocking state to prevent loops from occurring. The STA designates a single switch as the root bridge and uses it as the reference point for all path calculations. In the figure below, the root bridge (switch S1) is chosen through an election process. All switches participating in spanning tree exchange BPDU frames to determine which switch has the lowest bridge ID (BID) on the network. The switch with the lowest BID automatically becomes the root bridge for the STA calculations.

A BPDU is a messaging frame exchanged by switches for STP. Each BPDU contains a BID that identifies the switch that sent the BPDU. The BID contains a priority value, the MAC address of the sending switch, and an optional extended system ID. The lowest BID value is determined by the combination of these three fields.

After the root bridge has been determined, the STA calculates the shortest path to it. Each switch uses the STA to determine which ports to block. While the STA determines the best paths to the root bridge for all switch ports in the broadcast domain, traffic is prevented from being forwarded through the network. The STA considers both path and port costs when determining which ports to block. The path costs are calculated using port cost values associated with port speeds for each switch port along a given path. The sum of the port cost values determines the overall path cost to the root bridge. If there is more than one path to choose from, STA chooses the path with the lowest path cost.

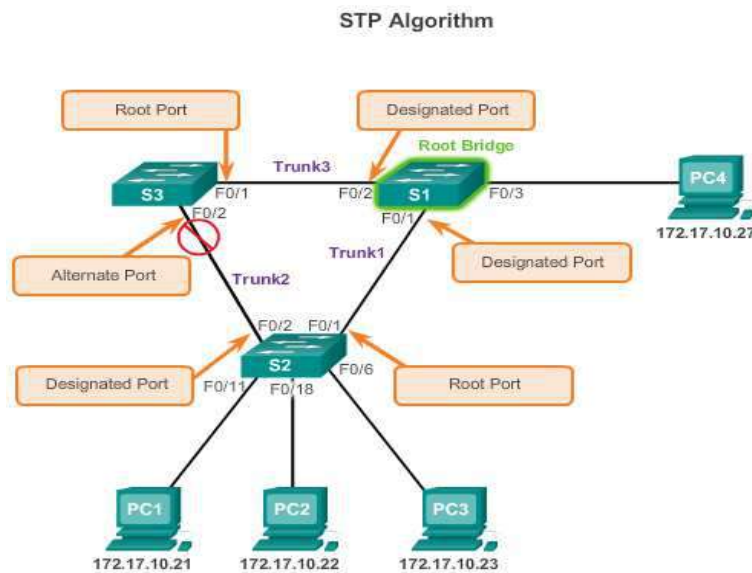
When the STA has determined which paths are most desirable relative to each switch, it assigns port roles to the participating switch ports. The port roles describe their relation in the network to the root bridge and whether they are allowed to forward traffic:

Root ports - Switch ports closest to the root bridge.

Designated ports - All non-root ports that are still permitted to forward traffic on the network.

Alternate and backup ports - Alternate ports and backup ports are configured to be in a blocking state to prevent loops. (Blocking ports only come into play when two ports on the same switch are connected to each other via a hub or single cable.)

Disabled ports - A disabled port is a switch port that is shut down.



Spanning Tree Algorithm – Root Bridge

Every spanning tree instance (switched LAN or broadcast domain) has a switch designated as the root bridge. The root bridge serves as a reference point for all spanning tree calculations to determine which redundant paths to block.

An election process determines which switch becomes the root bridge.

The BID is made up of a priority value, an extended system ID, and the MAC address of the switch.

All switches in the broadcast domain participate in the election process. After a switch boots, it begins to send out BPDUs every two seconds. These BPDUs contain the switch BID and the root ID.

As the switches forward their BPDU frames, adjacent switches in the broadcast domain read the root ID information from the BPDU frames. If the root ID from a BPDU received is lower than the root ID on the receiving switch, then the receiving switch updates its root ID, identifying the adjacent switch as the root bridge. Actually, it may not be an adjacent switch, but could be any other switch in the broadcast domain. The switch then forwards new BPDU frames with the lower root ID to the other adjacent switches. Eventually, the switch with the lowest BID ends up being identified as the root bridge for the spanning tree instance.

There is a root bridge elected for each spanning tree instance. It is possible to have multiple distinct root bridges. If all ports on all switches are members of VLAN 1, then there is only one spanning tree instance. The extended system ID plays a role in how spanning tree instances are determined.

Spanning Tree Algorithm – Path Cost

When the root bridge has been elected for the spanning tree instance, the STA starts the process of determining the best paths to the root bridge from all destinations in the broadcast domain. The path information is determined by summing up the individual port costs along the path from the destination to the root bridge. Each destination is actually a switch port.

The default port costs are defined by the speed at which the port operates.

Although switch ports have a default port cost associated with them, the port cost is configurable. The ability to configure individual port costs gives the administrator the flexibility to manually control the spanning tree paths to the root bridge.

To configure the port cost of an interface, enter the **spanning-tree cost value** command in interface configuration mode. The value can be between 1 and 200,000,000.

To restore the port cost back to the default value of 19, enter the **no spanning-tree cost** interface configuration mode command.

The path cost is equal to the sum of all the port costs along the path to the root bridge. Paths with the lowest cost become preferred, and all other redundant paths are blocked. To verify the port and path cost to the root bridge, enter the **show spanning-tree** command.

Topic 160: Spanning Tree Protocol (STP) – Part III:

802.1D BPDU Frame Format

The spanning tree algorithm depends on the exchange of BPDUs to determine a root bridge. A BPDU frame contains 12 distinct fields that convey path and priority information used to determine the root bridge and paths to the root bridge.

- The first four fields identify the protocol, version, message type, and status flags.
- The next four fields are used to identify the root bridge and the cost of the path to the root bridge.
- The last four fields are all timer fields that determine how frequently BPDU messages are sent and how long the information received through the BPDU process (next topic) is retained.

Field Number	Bytes	Field
1-4	2	Protocol ID
	1	Version
	1	Message type
	1	Flags
5-8	8	Root ID
	4	Cost of path
	8	Bridge ID
	2	Port ID
9-12	2	Message age
	2	Max age
	2	Hello time
	2	Forward delay

- The Protocol ID field indicates the type of protocol being used.
- The Version field indicates the version of the protocol.
- The Message type field indicates the type of message.
- The Flags field includes one of the following:
 - o Topology change (TC) bit, which signals a topology change in the event a path to the root bridge has been disrupted.

- Topology change acknowledgment (TCA) bit, which is set to acknowledge receipt of a configuration message with the TC bit set.
- The Root ID field indicates the root bridge by listing its 2-byte priority followed by its 6-byte MAC address ID. When a switch first boots, the root ID is the same as the bridge ID. However, as the election process occurs, the lowest bridge ID replaces the local root ID to identify the root bridge switch.
- The Cost of path field indicates the cost of the path from the bridge sending the configuration message to the root bridge.
- The Bridge ID field indicates the priority and MAC address ID of the bridge sending the message. This label allows the root bridge to identify where the BPDU originated, as well as for identifying the multiple paths from the switch to the root bridge. When the root bridge receives more than one BPDU from a switch with different path costs it knows that there are two distinct paths and uses the one path with the lower cost.
- The Port ID field indicates the port number from which the configuration message was sent. This field allows loops created by multiple attached bridges to be detected and corrected.
- The Message age field indicates the amount of time that has elapsed since the root sent the configuration message on which the current configuration message is based.
- The Max age field indicates when the current configuration message should be deleted. This is 20 seconds by default, but can be tuned to be between 6 and 40 seconds.
- The Hello time field indicates the time between root bridge configuration messages. The interval defines how long the root bridge waits between sending configuration message BPDUs. This is equal to 2 seconds by default, but can be tuned to be between 1 and 10 seconds.
- The Forward delay field indicates the length of time that bridges should wait before transitioning to a new state after a topology change. If a bridge transitions too soon, it is possible that not all network links will be ready to change their state and loops can result. This is by default equal to 15 seconds for each state, but can be tuned to be between 4 and 30 seconds.

Topic 161: Spanning Tree Protocol (STP) – Part IV:

BPDU Propagation and Process

Each switch in the broadcast domain initially assumes that it is the root bridge for a spanning tree instance, so the BPDU frames sent contain the BID of the local switch as the root ID. By default, BPDU frames are sent every two seconds after a switch is booted; that is, the default value of the Hello timer specified in the BPDU frame is two seconds. Each switch maintains local information about its own BID, the root ID, and the path cost to the root.

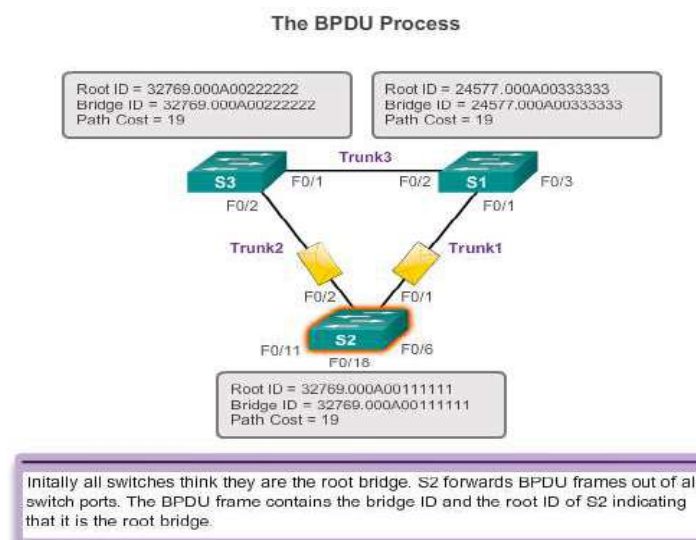
When adjacent switches receive a BPDU frame, they compare the root ID from the BPDU frame with the local root ID. If the root ID in the BPDU is lower than the local root ID, the switch updates the local root ID and the ID in its BPDU messages. These messages indicate the new root bridge on the network. The distance to the root bridge is also indicated by the path cost update. For example, if the BPDU was received on a Fast Ethernet switch port, the path cost would increment by 19. If the local root ID is lower than the root ID received in the BPDU frame, the BPDU frame is discarded.

After a root ID has been updated to identify a new root bridge, all subsequent BPDU frames sent from that switch contain the new root ID and updated path cost. That way, all other adjacent switches are able to see the lowest root ID identified at all times. As the BPDU frames pass between other adjacent switches, the path cost is continually updated to indicate the total path cost to the root bridge. Each switch in the spanning tree uses its path costs to identify the best possible path to the root bridge.

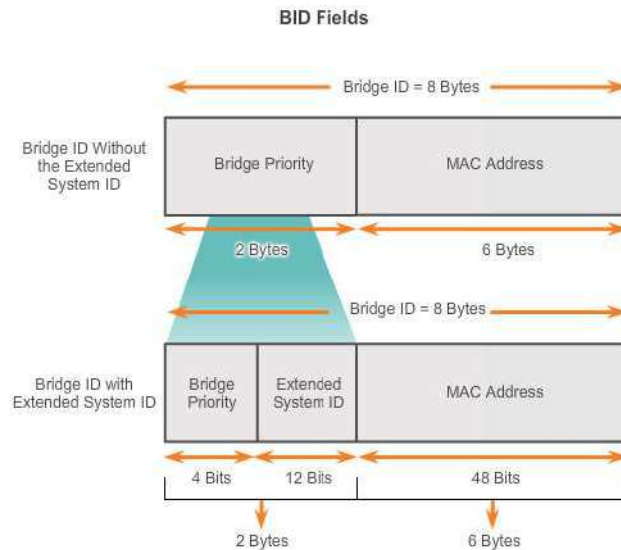
The following summarizes the BPDU process:

Note: Priority is the initial deciding factor when electing a root bridge. If the priorities of all the switches are the same, the device with the lowest MAC address becomes the root bridge.

1. Initially, each switch identifies itself as the root bridge. S2 forwards BPDU frames out all switch ports.
2. When S3 receives a BPDU from switch S2, S3 compares its root ID with the BPDU frame it received. The priorities are equal, so the switch is forced to examine the MAC address portion to determine which MAC address has a lower value. Because S2 has a lower MAC address value, S3 updates its root ID with the S2 root ID. At that point, S3 considers S2 as the root bridge.
3. When S1 compares its root ID with the one in the received BPDU frame, it identifies its local root ID as the lower value and discards the BPDU from S2.
4. When S3 sends out its BPDU frames, the root ID contained in the BPDU frame is that of S2.
5. When S2 receives the BPDU frame, it discards it after verifying that the root ID in the BPDU matched its local root ID.
6. Because S1 has a lower priority value in its root ID, it discards the BPDU frame received from S3.
7. S1 sends out its BPDU frames.
8. S3 identifies the root ID in the BPDU frame as having a lower value and, therefore, updates its root ID values to indicate that S1 is now the root bridge.
9. S2 identifies the root ID in the BPDU frame as having a lower value and, therefore, updates its root ID values to indicate that S1 is now the root bridge.



Extended System ID



The bridge ID (BID) is used to determine the root bridge on a network. The BID field of a BPDU frame contains three separate fields:

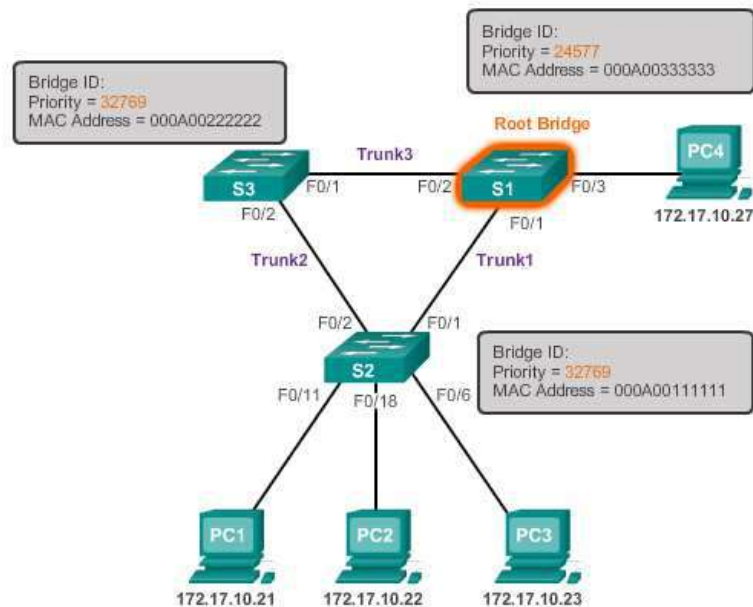
- Bridge priority
- Extended system ID
- MAC address

Each field is used during the root bridge election.

Bridge Priority

The bridge priority is a customizable value that can be used to influence which switch becomes the root bridge. The switch with the lowest priority, which implies the lowest BID, becomes the root bridge because a lower priority value takes precedence. For example, to ensure that a specific switch is always the root bridge, set the priority to a lower value than the rest of the switches on the network. The default priority value for all Cisco switches is 32768. The range is 0 to 61440 in increments of 4096. Valid priority values are 0, 4096, 8192, 12288, 16384, 20480, 24576, 28672, 32768, 36864, 40960, 45056, 49152, 53248, 57344, and 61440. All other values are rejected. A bridge priority of 0 takes precedence over all other bridge priorities.

Priority-based decision



Extended System ID

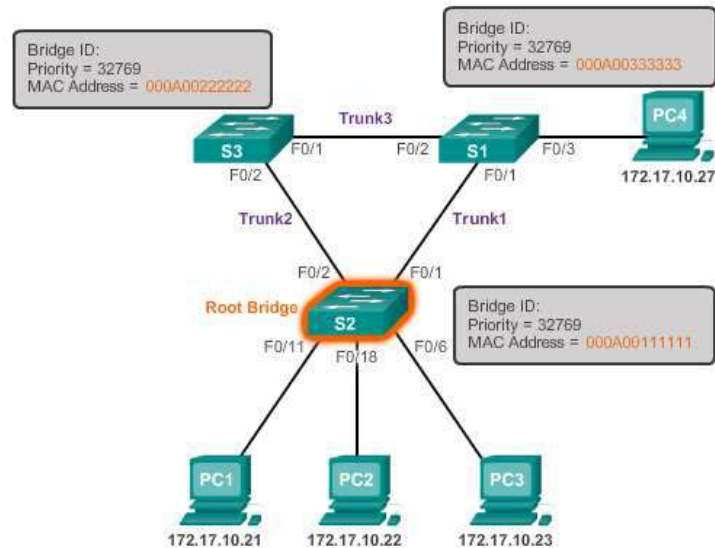
Early implementations of IEEE 802.1D were designed for networks that did not use VLANs. There was a single common spanning tree across all switches. For this reason, in older Cisco switches, the extended system ID could be omitted in BPDU frames. As VLANs became common for network infrastructure segmentation, 802.1D was enhanced to include support for VLANs, requiring the VLAN ID to be included in the BPDU frame. VLAN information is included in the BPDU frame through the use of the extended system ID. All newer switches include the use of the extended system ID by default.

The extended system ID value is added to the bridge priority value in the BID to identify the priority and VLAN of the BPDU frame.

When two switches are configured with the same priority and have the same extended system ID, the switch having the MAC address with the lowest hexadecimal value will have the lower BID. Initially, all switches are configured with the same default priority value. The MAC address is then the deciding factor on which switch is going to become the root bridge. To ensure that the root bridge decision best meets network requirements, it is recommended that the administrator configure the desired root bridge switch with a lower priority. This also ensures that the addition of new switches to the network does not trigger a new spanning tree election, which can disrupt network communication while a new root bridge is being selected.

In Figure above, S1 has a lower priority than the other switches; therefore, it is preferred as the root bridge for that spanning tree instance.

When all switches are configured with the same priority, as is the case with all switches kept in the default configuration with a priority of 32768, the MAC address becomes the deciding factor for which switch becomes the root bridge.



The MAC address with the lowest hexadecimal value is considered to be the preferred root bridge. In the example above, S2 has the lowest value for its MAC address and is, therefore, designated as the root bridge for that spanning tree instance.

Topic 162: Varieties of Spanning Tree Protocols:

Several varieties of spanning tree protocols have emerged after the original IEEE 802.1D.

The varieties of spanning tree protocols include:

STP - This is the original IEEE 802.1D version (802.1D-1998 and earlier) that provides a loop-free topology in a network with redundant links. Common Spanning Tree (CST) assumes one spanning tree instance for the entire bridged network, regardless of the number of VLANs.

PVST+ - This is a Cisco enhancement of the original 802.1D standard that provides a separate 802.1D spanning tree instance for each VLAN configured in the network. The separate instance supports PortFast, UplinkFast, BackboneFast, BPDU guard, BPDU filter, root guard, and loop guard.

802.1D-2004 - This is an updated version of the STP standard, incorporating IEEE 802.1w.

Rapid Spanning Tree Protocol (RSTP) or IEEE 802.1w - This is an evolution of STP that provides faster convergence than STP.

Rapid PVST+ - This is a Cisco enhancement of RSTP that uses PVST+. Rapid PVST+ provides a separate instance of 802.1w per VLAN. The separate instance supports PortFast, BPDU guard, BPDU filter, root guard, and loop guard.

Multiple Spanning Tree Protocol(MSTP) - This is an IEEE standard inspired by the earlier Cisco proprietary Multiple Instance STP (MISTP) implementation. MSTP maps multiple VLANs into the same spanning tree instance. The Cisco implementation of MSTP is MST, which provides up to 16 instances of RSTP and combines many VLANs with the same physical and logical topology into a common RSTP instance. Each instance supports PortFast, BPDU guard, BPDU filter, root guard, and loop guard.

A network professional, whose duties include switch administration, may be required to decide which type of spanning tree protocol to implement.

To analyze the STP topology, follow these steps:

Step 1. Discover the Layer 2 topology. Use network documentation if it exists or use the **show cdp neighbors** command to discover the Layer 2 topology.

Step 2. After discovering the Layer 2 topology, use STP knowledge to determine the expected Layer 2 path. It is necessary to know which switch is the root bridge.

Step 3. Use the **show spanning-tree vlan** command to determine which switch is the root bridge.

Step 4. Use the **show spanning-tree vlan** command on all switches to find out which ports are in blocking or forwarding state and confirm your expected Layer 2 path.

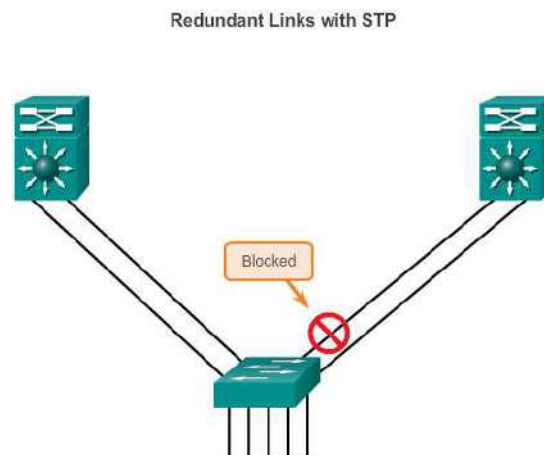
Topic 163: Link Aggregation:

In the figure below, traffic coming from several links (usually 100 or 1000 Mb/s) aggregates on the access switch and must be sent to distribution switches. Because of the traffic aggregation, links with higher bandwidth must be available between the access and distribution switches.

It may be possible to use faster links, such as 10 Gb/s, on the aggregated link between the access and distribution layer switches. However, adding faster links is expensive. Additionally, as the speed increases on the access links, even the fastest possible port on the aggregated link is no longer fast enough to aggregate the traffic coming from all access links.

It is also possible to multiply the number of physical links between the switches to increase the overall speed of switch-to-switch communication. However, by default, STP is enabled on switch devices. STP will block redundant links to prevent routing loops.

For these reasons, the best solution is to implement an EtherChannel configuration.



By default, STP will block redundant links.

EtherChannel

EtherChannel technology was originally developed by Cisco as a LAN switch-to-switch technique of grouping several Fast Ethernet or Gigabit Ethernet ports into one logical channel. When an EtherChannel is configured, the resulting virtual interface is called a port channel. The physical interfaces are bundled together into a port channel interface.

EtherChannel technology has many advantages:

- Most configuration tasks can be done on the EtherChannel interface instead of on each individual port, ensuring configuration consistency throughout the links.
- EtherChannel relies on existing switch ports. There is no need to upgrade the link to a faster and more expensive connection to have more bandwidth.
- Load balancing takes place between links that are part of the same EtherChannel. Depending on the hardware platform, one or more load-balancing methods can be implemented. These methods include source MAC to destination MAC load balancing, or source IP to destination IP load balancing, across the physical links.
- EtherChannel creates an aggregation that is seen as one logical link. When several EtherChannel bundles exist between two switches, STP may block one of the bundles to prevent switching loops. When STP blocks one of the redundant links, it blocks the entire EtherChannel. This blocks all the ports belonging to that EtherChannel link. Where there is only one EtherChannel link, all physical links in the EtherChannel are active because STP sees only one (logical) link.
- EtherChannel provides redundancy because the overall link is seen as one logical connection. Additionally, the loss of one physical link within the channel does not create a change in the topology; therefore a spanning tree recalculation is not required. Assuming at least one physical link is present; the EtherChannel remains functional, even if its overall throughput decreases because of a lost link within the EtherChannel.
- EtherChannel can be implemented by grouping multiple physical ports into one or more logical EtherChannel links.

The original purpose of EtherChannel is to increase speed capability on aggregated links between switches. However, this concept was extended as EtherChannel technology became more popular, and now many servers also support link aggregation with EtherChannel. EtherChannel creates a one-to-one relationship; that is, one EtherChannel link connects only two devices. An EtherChannel link can be created between two switches or an EtherChannel link can be created between an EtherChannel-enabled server and a switch. However, traffic cannot be sent to two different switches through the same EtherChannel link.

The individual EtherChannel group member port configuration must be consistent on both devices. If the physical ports of one side are configured as trunks, the physical ports of the other side must also be configured as trunks within the same native VLAN. Additionally, all ports in each EtherChannel link must be configured as Layer 2 ports.

PAgP

PAgP is a Cisco-proprietary protocol that aids in the automatic creation of EtherChannel links. When an EtherChannel link is configured using PAgP, PAgP packets are sent between EtherChannel-capable ports to negotiate the forming of a channel.

Link Aggregation Control Protocol (LACP)

LACP is part of an IEEE specification (802.3ad) that allows several physical ports to be bundled to form a single logical channel. LACP allows a switch to negotiate an automatic bundle by sending LACP packets to the peer. It performs a function similar to PAgP with Cisco EtherChannel. Because LACP is an IEEE standard, it can be used to facilitate EtherChannels in multivendor environments. On Cisco devices, both protocols are supported.

Topic 164: Packet Tracer–Switched Network with Redundant Links:

- Build the Network
- Determine the Root Bridge

Topic 165: Inter-VLAN Routing:

VLANs are used to segment switched networks. Layer 2 switches, such as the Catalyst 2960 Series, can be configured by a network professional with over 4,000 VLANs. However, Layer 2 switches have very limited IPv4 and IPv6 functionality and cannot perform the routing function of routers. While Layer 2 switches are gaining more IP functionality, such as the ability to perform static routing, these switches do not support dynamic routing. With the large number of VLANs possible on these switches, static routing is insufficient.

A VLAN is a broadcast domain, so computers on separate VLANs are unable to communicate without the intervention of a routing device. Any device that supports Layer 3 routing, such as a router or a multilayer switch, can be used to perform the necessary routing functionality. Regardless of the device used, the process of forwarding network traffic from one VLAN to another VLAN using routing is known as inter-VLAN routing.

Legacy Inter-VLAN Routing

Historically, the first solution for inter-VLAN routing relied on routers with multiple physical interfaces. Each interface had to be connected to a separate network and configured with a distinct subnet.

In this legacy approach, inter-VLAN routing is performed by connecting different physical router interfaces to different physical switch ports. The switch ports connected to the router are placed in access mode and each physical interface is assigned to a different VLAN. Each router interface can then accept traffic from the VLAN associated with the switch interface that it is connected to, and traffic can be routed to the other VLANs connected to the other interfaces.

Router-on-a-Stick Inter-VLAN Routing

While legacy inter-VLAN routing requires multiple physical interfaces on both the router and the switch, a more common, present-day implementation of inter-VLAN routing does not. Instead, some router software permits configuring a router interface as a trunk link, meaning only one physical interface is required on the router and the switch to route packets between multiple VLANs.

Router-on-a-stick is a type of router configuration in which a single physical interface routes traffic between multiple VLANs on a network. The router interface is configured to operate as a trunk link and is connected to a switch port that is configured in trunk mode. The router performs inter-VLAN routing by accepting VLAN-tagged traffic on the trunk interface coming from the adjacent switch, and then internally routing between the VLANs using subinterfaces. The router then forwards the routed traffic, VLAN-tagged for the destination VLAN, out the same physical interface as it used to receive the traffic.

Subinterfaces are software-based virtual interfaces, associated with a single physical interface. Subinterfaces are configured in software on a router and each subinterface is independently configured with an IP address and VLAN assignment. Subinterfaces are configured for different subnets corresponding to their VLAN assignment to facilitate logical routing. After a routing decision is made based on the destination VLAN, the data frames are VLAN-tagged and sent back out the physical interface.

Multi-Layer Switch Inter-VLAN Routing

The router-on-a-stick implementation of inter-VLAN routing requires only one physical interface on a router and one interface on a switch, simplifying the cabling of the router. However, in other implementations of inter-VLAN routing, a dedicated router is not required.

Multilayer switches can perform Layer 2 and Layer 3 functions, replacing the need for dedicated routers to perform basic routing on a network. Multilayer switches support dynamic routing and inter-VLAN routing.

To enable a multilayer switch to perform routing functions, the multilayer switch must have IP routing enabled.

Multilayer switching is more scalable than any other inter-VLAN routing implementation. This is because routers have a limited number of available ports to connect to networks. Additionally, for interfaces that are configured as a trunk line, limited amounts of traffic can be accommodated on that line at one time.

With a multilayer switch, traffic is routed internal to the switch device, which means packets are not filtered down a single trunk line to obtain new VLAN-tagging information. A multilayer switch does not, however, completely replace the functionality of a router. Routers support a significant number of additional features, such as the ability to implement greater security controls. Rather, a multilayer switch can be thought of as a Layer 2 device that is upgraded to have some routing capabilities.

Topic 166: Configure Legacy Inter-VLAN Routing:

Legacy inter-VLAN routing requires routers to have multiple physical interfaces. The router accomplishes the routing by having each of its physical interfaces connected to a unique VLAN. Each interface is also configured with an IP address for the subnet associated with the particular VLAN to which it is connected. By configuring the IP addresses on the physical interfaces, network devices connected to each of the VLANs can communicate with the router using the physical interface connected to the same VLAN. In this configuration, network devices can use the router as a gateway to access the devices connected to the other VLANs.

The routing process requires the source device to determine if the destination device is local or remote to the local subnet. The source device accomplishes this by comparing the source and destination IP addresses against the subnet mask. When the destination IP address has been determined to be on a remote network, the source device must identify where it needs to forward the packet to reach the destination device. The source device examines the local routing table to determine where it needs to send the data. Devices use their default gateway as the Layer 2 destination for all traffic that must leave the local subnet. The default gateway is the route that the device uses when it has no other explicitly defined route to the destination network. The IP address of the router interface on the local subnet acts as the default gateway for the sending device.

When the source device has determined that the packet must travel through the local router interface on the connected VLAN, the source device sends out an ARP request to determine the MAC address of the local router interface. When the router sends its ARP reply back to the source device, the source device can use the MAC address to finish framing the packet before it sends it out on the network as unicast traffic.

Because the Ethernet frame has the destination MAC address of the router interface, the switch knows exactly which switch port to forward the unicast traffic out of to reach the router interface for that VLAN. When the frame arrives at the router, the router removes the source and destination MAC address information to examine the destination IP address of the packet. The router compares the destination address to entries in its routing table to determine where it needs to forward the data to reach its final destination. If the router determines that the destination network is a locally connected network, as is the case with inter-VLAN routing, the router sends an ARP request out the interface physically connected to the destination VLAN. The destination device responds back to the router with its MAC address, which the router then uses to frame the packet. The router then sends the unicast traffic to the switch, which forwards it out the port where the destination device is connected.

Legacy Inter-VLAN Routing-Switch Configuration

To configure legacy inter-VLAN routing, start by configuring the switch.

Use the **vlan *vlan_id*** global configuration mode command to create VLANs. After the VLANs have been created, the switch ports are assigned to the appropriate VLANs. The **switchport access vlan *vlan_id*** command is executed from interface configuration mode on the switch for each interface to which the router connects.

Finally, to protect the configuration so that it is not lost after a reload of the switch, the **copy running-config startup-config** command is executed to back up the running configuration to the startup configuration.

Next, the router can be configured to perform inter-VLAN routing.

Router interfaces are configured in a manner similar to configuring VLAN interfaces on switches. To configure a specific interface, change to interface configuration mode from global configuration mode.

Router interfaces are disabled by default and must be enabled using the **no shutdown** command before they are used. After the **no shutdown** interface configuration mode command has been issued, a

notification displays, indicating that the interface state has changed to up. This indicates that the interface is now enabled.

The process is repeated for all router interfaces. Each router interface must be assigned to a unique subnet for routing to occur.

Topic 167: Configure Router-on-a-Stick Inter-VLAN Routing:

Legacy inter-VLAN routing using physical interfaces has a significant limitation. Routers have a limited number of physical interfaces to connect to different VLANs. As the number of VLANs increases on a network, having one physical router interface per VLAN quickly exhausts the physical interface capacity of a router. An alternative in larger networks is to use VLAN trunking and subinterfaces. VLAN trunking allows a single physical router interface to route traffic for multiple VLANs. This technique is termed router-on-a-stick and uses virtual subinterfaces on the router to overcome the hardware limitations based on physical router interfaces.

Subinterfaces are software-based virtual interfaces that are assigned to physical interfaces. Each subinterface is configured independently with its own IP address and subnet mask. This allows a single physical interface to simultaneously be part of multiple logical networks.

When configuring inter-VLAN routing using the router-on-a-stick model, the physical interface of the router must be connected to a trunk link on the adjacent switch. On the router, subinterfaces are created for each unique VLAN on the network. Each subinterface is assigned an IP address specific to its subnet/VLAN and is also configured to tag frames for that VLAN. This way, the router can keep the traffic from each subinterface separated as it traverses the trunk link back to the switch.

Functionally, the router-on-a-stick model is the same as using the legacy inter-VLAN routing model, but instead of using the physical interfaces to perform the routing, subinterfaces of a single physical interface are used.

Topic 168: Layer 3 Switching:

Router-on-a-stick is simple to implement because routers are usually available in every network. Most enterprise networks use multilayer switches to achieve high-packet processing rates using hardware-based switching. Layer 3 switches usually have packet-switching throughputs in the millions of packets per second (pps), whereas traditional routers provide packet switching in the range of 100,000 pps to more than 1 million pps.

All Catalyst multilayer switches support the following types of Layer 3 interfaces:

Routed port - A pure Layer 3 interface similar to a physical interface on a Cisco IOS router.

Switch virtual interface (SVI) - A virtual VLAN interface for inter-VLAN routing. In other words, SVIs are the virtual-routed VLAN interfaces.

High-performance switches, such as the Catalyst 6500 and Catalyst 4500, perform almost every function involving OSI Layer 3 and higher using hardware-based switching that is based on Cisco Express Forwarding.

All Layer 3 Cisco Catalyst switches support routing protocols, but several models of Catalyst switches require enhanced software for specific routing protocol features. Catalyst 2960 Series switches running IOS Release 12.2(55) or later, support static routing.

Catalyst switches use different default settings for interfaces. All members of the Catalyst 3560 and 4500 families of switches use Layer 2 interfaces by default. Members of the Catalyst 6500 family of switches running Cisco IOS use Layer 3 interfaces by default. Default interface configurations do not appear in the running or startup configuration. Depending on which Catalyst family of switches is used, the **switchport** or **no switchport** interface configuration mode commands might be present in the running config or startup configuration files.

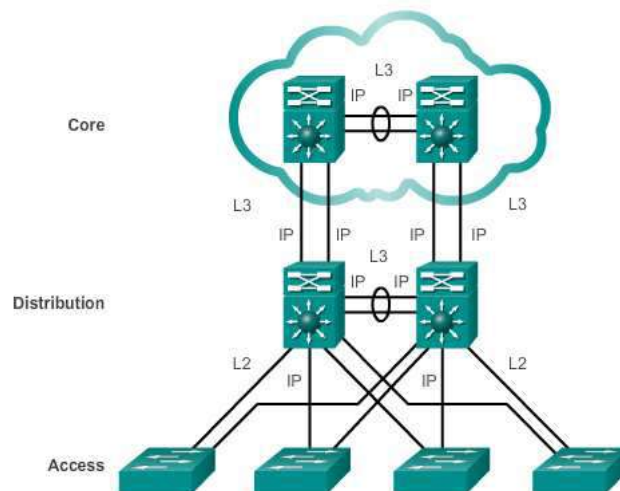
Inter-VLAN Routing with Switch Virtual Interfaces

In the early days of switched networks, switching was fast (often at hardware speed, meaning the speed was equivalent to the time it took to physically receive and forward frames onto other ports) and routing was slow (routing had to be processed in software). This prompted network designers to extend the switched portion of the network as much as possible. Access, distribution, and core layers were often configured to communicate at Layer 2. This topology created loop issues. To solve these issues, spanning-tree technologies were used to prevent loops while still enabling flexibility and redundancy in inter-switch connections.

However, as network technologies have evolved, routing has become faster and cheaper. Today, routing can be performed at wire speed. One consequence of this evolution is that routing can be transferred to the core and the distribution layers without impacting network performance.

Many users are in separate VLANs, and each VLAN is usually a separate subnet. Therefore, it is logical to configure the distribution switches as Layer 3 gateways for the users of each access switch VLAN. This implies that each distribution switch must have IP addresses matching each access switch VLAN.

As shown in the figure below, Layer 3 (routed) ports are normally implemented between the distribution and the core layer. The network architecture depicted is not dependent on spanning tree because there are no physical loops in the Layer 2 portion of the topology.



Switch Virtual Interfaces

An SVI is a virtual interface that is configured within a multilayer switch. An SVI can be created for any VLAN that exists on the switch. An SVI is considered to be virtual because there is no physical port dedicated to the interface. It can perform the same functions for the VLAN as a router interface would, and can be configured in much the same way as a router interface (i.e., IP address, inbound/outbound ACLs, etc.). The SVI for the VLAN provides Layer 3 processing for packets to or from all switch ports associated with that VLAN.

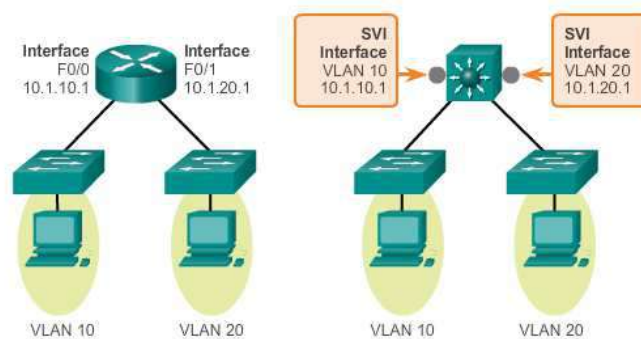
Whenever the SVI is created, ensure that particular VLAN is present in the VLAN database. In the figure below, the switch should have VLAN 10 and VLAN 20 present in the VLAN database; otherwise, the SVI interface stays down.

The following are some of the reasons to configure SVI:

- To provide a gateway for a VLAN so that traffic can be routed into or out of that VLAN
- To provide Layer 3 IP connectivity to the switch
- To support routing protocol and bridging configurations

The following are some of the advantages of SVIs (the only disadvantage is that multilayer switches are more expensive):

- It is much faster than router-on-a-stick, because everything is hardware switched and routed.
- No need for external links from the switch to the router for routing.
- Not limited to one link. Layer 2 EtherChannels can be used between the switches to get more bandwidth.
- Latency is much lower, because it does not need to leave the switch.



Inter-VLAN Routing with Routed Ports

Routed Ports and Access Ports on a Switch

A routed port is a physical port that acts similarly to an interface on a router. Unlike an access port, a routed port is not associated with a particular VLAN. A routed port behaves like a regular router interface. Also, because Layer 2 functionality has been removed, Layer 2 protocols, such as STP, do not function on a routed interface. However, some protocols, such as LACP and EtherChannel, do function at Layer 3.

Unlike Cisco IOS routers, routed ports on a Cisco IOS switch do not support subinterfaces.

Routed ports are used for point-to-point links. Connecting WAN routers and security devices are examples of the use of routed ports. In a switched network, routed ports are mostly configured between switches in the core and distribution layer. The figure below illustrates an example of routed ports in a campus switched network.

To configure routed ports, use the **no switchport** interface configuration mode command on the appropriate ports. For example, the default configuration of the interfaces on Catalyst 3560 switches are Layer 2 interfaces, so they must be manually configured as routed ports. In addition, assign an IP address and other Layer 3 parameters as necessary. After assigning the IP address, verify that IP routing is globally enabled and that applicable routing protocols are configured.

