



Configuring IP Unicast Routing

This chapter describes how to configure IP unicast routing on your Catalyst 3550 multilayer switch. Beginning with Cisco IOS Release 12.1(11)EA1, basic routing functions, including static unicast routing and the Routing Information Protocol (RIP), are available with both the IP base image (formerly known as the standard multilayer software image [SMI]) and the IP services image (formerly known as the enhanced multilayer software image [EMI]). To use advanced routing features and other routing protocols, or for all routing support prior to Cisco IOS Release 12.1(11)EA1, you must have the IP services image installed on your switch.

For more detailed IP unicast configuration information, see the *Cisco IOS IP Configuration Guide, Release 12.2*. For complete syntax and usage information for the commands used in this chapter, see these command references:

- *Cisco IOS IP Command Reference, Volume 1 of 3: Addressing and Services, Release 12.2*
- *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*
- *Cisco IOS IP Command Reference, Volume 3 of 3: Multicast, Release 12.2*

This chapter consists of these sections:

- [Understanding IP Routing, page 32-2](#)
- [Steps for Configuring Routing, page 32-3](#)
- [Configuring IP Addressing on Layer 3 Interfaces, page 32-4](#)
- [Enabling IP Unicast Routing, page 32-18](#)
- [Configuring RIP, page 32-19](#)
- [Configuring OSPF, page 32-24](#)
- [Configuring EIGRP, page 32-34](#)
- [Configuring BGP, page 32-41](#)
- [Configuring Multi-VRF CE, page 32-61](#)
- [Configuring Protocol-Independent Features, page 32-72](#)
- [Monitoring and Maintaining the IP Network, page 32-85](#)



Note

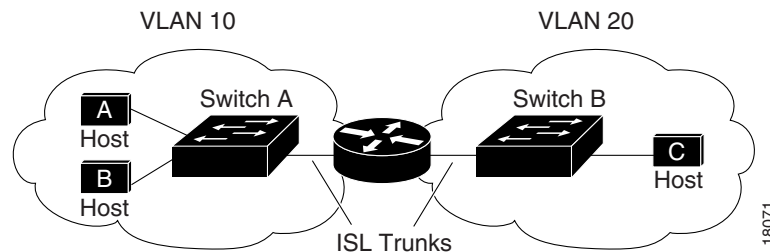
When configuring routing parameters on the switch, to allocate system resources to maximize the number of unicast routes allowed, you can use the **sdm prefer routing** global configuration command to set the Switch Database Management (sdm) feature to the routing template. For more information on the SDM templates, see the [“Optimizing System Resources for User-Selected Features” section on page 6-26](#).

Understanding IP Routing

In some network environments, VLANs are associated with individual networks or subnetworks. In an IP network, each subnetwork is mapped to an individual VLAN. Configuring VLANs helps control the size of the broadcast domain and keeps local traffic local. However, network devices in different VLANs cannot communicate with one another without a Layer 3 device (router) to route traffic between the VLAN, referred to as inter-VLAN routing. You configure one or more routers to route traffic to the appropriate destination VLAN.

Figure 32-1 shows a basic routing topology. Switch A is in VLAN 10, and Switch B is in VLAN 20. The router has an interface in each VLAN.

Figure 32-1 Routing Topology Example



When Host A in VLAN 10 needs to communicate with Host B in VLAN 10, it sends a packet addressed to that host. Switch A forwards the packet directly to Host B, without sending it to the router.

When Host A sends a packet to Host C in VLAN 20, Switch A forwards the packet to the router, which receives the traffic on the VLAN 10 interface. The router checks the routing table, determines the correct outgoing interface, and forwards the packet on the VLAN 20 interface to Switch B. Switch B receives the packet and forwards it to Host C.

Routers can perform unicast routing in three different ways:

- By using default routing
- By using preprogrammed static routes for the traffic
- By dynamically calculating routes by using a routing protocol

Default routing refers to sending traffic with a destination unknown to the router to a default outlet or destination.

Static unicast routing forwards packets from predetermined ports through a single path into and out of a network. Static routing is secure and uses little bandwidth, but does not automatically respond to changes in the network, such as link failures, and therefore, might result in unreachable destinations. As networks grow, static routing becomes a labor-intensive liability.

Dynamic routing protocols are used by routers to dynamically calculate the best route for forwarding traffic. There are two types of dynamic routing protocols:

- Routers using distance-vector protocols maintain routing tables with distance values of networked resources, and periodically pass these tables to their neighbors. Distance-vector protocols use one or a series of metrics for calculating the best routes. These protocols are easy to configure and use.
- Routers using link-state protocols maintain a complex database of network topology, based on the exchange of link-state advertisements (LSAs) between routers. LSAs are triggered by an event in the network, which speeds up the convergence time or time required to respond to these changes. Link-state protocols respond quickly to topology changes, but require greater bandwidth and more resources than distance-vector protocols.

Distance-vector protocols supported by the switch are Routing Information Protocol (RIP), which uses a single distance metric (cost) to determine the best path and Border Gateway Protocol (BGP), which adds a path vector mechanism. The switch also supports the Open Shortest Path First (OSPF) link-state protocol and Enhanced Interior Gateway Routing Protocol (EIGRP), which adds some link-state routing features to traditional IGRP to improve efficiency.

**Note**

The IP base image supports only default routing, static routing, and RIP. All other routing protocols require the IP services image on your switch.

Steps for Configuring Routing

By default, IP routing is disabled on the switch, and you must enable it before routing can take place. For detailed IP routing configuration information, see the *Cisco IOS IP Configuration Guide, Release 12.2*.

In the following procedures, the specified interface must be one of these Layer 3 interfaces:

- A routed port: a physical port configured as a Layer 3 port by using the **no switchport** interface configuration command.
- A switch virtual interface (SVI): a VLAN interface created by using the **interface vlan** *vlan_id* global configuration command and by default a Layer 3 interface.
- An EtherChannel port channel in Layer 3 mode: a port-channel logical interface created by using the **interface port-channel** *port-channel-number* global configuration command and binding the Ethernet interface into the channel group. For more information, see the [“Configuring Layer 3 EtherChannels” section on page 31-12](#).

**Note**

The switch does not support tunnel interfaces for unicast routed traffic.

All Layer 3 interfaces must have IP addresses assigned to them. See the [“Assigning IP Addresses to Network Interfaces” section on page 32-5](#).

**Note**

A Layer 3 switch can have an IP address assigned to each routed port and SVI. The number of routed ports and SVIs that you can configure is not limited by software. However, the interrelationship between this number and the number and volume of features being implemented might have an impact on CPU utilization because of hardware limitations. For more information about feature combinations, see the [“Optimizing System Resources for User-Selected Features” section on page 6-26](#).

Configuring routing consists of several main procedures:

- To support VLAN interfaces, create and configure VLANs on the switch, and assign VLAN membership to Layer 2 interfaces. For more information, see [Chapter 11, “Configuring VLANs.”](#)
- Configure Layer 3 interfaces.
- Enable IP routing on the switch.
- Assign IP addresses to the Layer 3 interfaces.
- Enable selected routing protocols on the switch.
- Configure routing protocol parameters (optional).

Configuring IP Addressing on Layer 3 Interfaces

A required task for configuring IP routing is to assign IP addresses to Layer 3 network interfaces to enable the interfaces and allow communication with the hosts on those interfaces that use IP. These sections describe how to configure various IP addressing features. Assigning IP addresses to the interface is required; the other procedures are optional.

- [Default Addressing Configuration, page 32-4](#)
- [Assigning IP Addresses to Network Interfaces, page 32-5](#)
- [Configuring Address Resolution Methods, page 32-8](#)
- [Routing Assistance When IP Routing is Disabled, page 32-11](#)
- [Configuring Broadcast Packet Handling, page 32-13](#)
- [Monitoring and Maintaining IP Addressing, page 32-17](#)

Default Addressing Configuration

[Table 32-1](#) shows the default addressing configuration.

Table 32-1 **Default Addressing Configuration**

Feature	Default Setting
IP address	None defined.
ARP	No permanent entries in the Address Resolution Protocol (ARP) cache. Encapsulation: Standard Ethernet-style ARP. Timeout: 14400 seconds (4 hours).
IP broadcast address	255.255.255.255 (all ones).
IP classless routing	Enabled.
IP default gateway	Disabled.
IP directed broadcast	Disabled (all IP directed broadcasts are dropped).
IP domain	Domain list: No domain names defined. Domain lookup: Enabled. Domain name: Enabled.
IP forward-protocol	If a helper address is defined or User Datagram Protocol (UDP) flooding is configured, UDP forwarding is enabled on default ports. Any-local-broadcast: Disabled. Spanning Tree Protocol (STP): Disabled. Turbo-flood: Disabled.
IP helper address	Disabled.
IP host	Disabled.

Table 32-1 *Default Addressing Configuration (continued)*

Feature	Default Setting
IRDP	Disabled. Defaults when enabled: - Broadcast IRDP advertisements. - Maximum interval between advertisements: 600 seconds. - Minimum interval between advertisements: 0.75 times max interval - Preference: 0.
IP proxy ARP	Enabled.
IP routing	Disabled.
IP subnet-zero	Disabled.

Assigning IP Addresses to Network Interfaces

An IP address identifies a location to which IP packets can be sent. Some IP addresses are reserved for special uses and cannot be used for host, subnet, or network addresses. [Table 32-2](#) lists ranges of IP addresses and shows which are reserved and which are available for use. RFC 1166, “Internet Numbers,” contains the official description of IP addresses.

Table 32-2 *Reserved and Available IP Addresses*

Class	Address or Range	Status
A	0.0.0.0	Reserved
	1.0.0.0 to 126.0.0.0	Available
	127.0.0.0	Reserved
B	128.0.0.0 to 191.254.0.0	Available
	191.255.0.0	Reserved
C	192.0.0.0	Reserved
	192.0.1.0 to 223.255.254	Available
	223.255.255.0	Reserved
D	224.0.0.0 to 239.255.255.255	Multicast group addresses
E	240.0.0.0 to 255.255.255.254	Reserved
	255.255.255.255	Broadcast

An interface can have one primary IP address. A mask identifies the bits that denote the network number in an IP address. When you use the mask to subnet a network, the mask is referred to as a subnet mask. To receive an assigned network number, contact your Internet service provider.

Beginning in privileged EXEC mode, follow these steps to assign an IP address and a network mask to a Layer 3 interface:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the Layer 3 interface to configure.
Step 3	no switchport	Remove the interface from Layer 2 configuration mode (if it is a physical interface).
Step 4	ip address <i>ip-address subnet-mask</i>	Configure the IP address and IP subnet mask.
Step 5	no shutdown	Enable the interface.
Step 6	end	Return to privileged EXEC mode.
Step 7	show interfaces [<i>interface-id</i>] show ip interface [<i>interface-id</i>] show running-config interface [<i>interface-id</i>]	Verify your entries.
Step 8	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no ip address** interface configuration command to remove an IP address or to disable IP processing.

This example shows how to configure an IP address on and enable Gigabit Ethernet interface 0/10:

```
Switch# configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)# interface gigabitethernet0/10
Switch(config-if)# no switchport
Switch(config-if)# ip address 10.1.2.3 255.255.0.0
Switch(config-if)# no shutdown
```

Use of Subnet Zero

Subnetting with a subnet address of zero is strongly discouraged because of the problems that can arise if a network and a subnet have the same addresses. For example, if network 131.108.0.0 is subnetted as 255.255.255.0, subnet zero would be written as 131.108.0.0, which is the same as the network address.

You can use the all ones subnet (131.108.255.0) and even though it is discouraged, you can enable the use of subnet zero if you need the entire subnet space for your IP address.

Beginning in privileged EXEC mode, follow these steps to enable subnet zero:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip subnet-zero	Enable the use of subnet zero for interface addresses and routing updates.
Step 3	end	Return to privileged EXEC mode.
Step 4	show running-config	Verify your entry.
Step 5	copy running-config startup-config	(Optional) Save your entry in the configuration file.

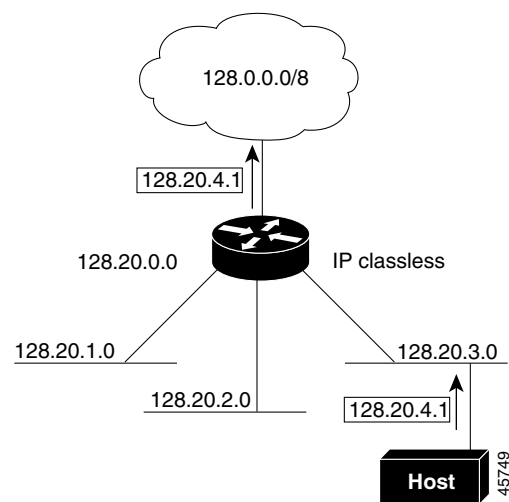
Use the **no ip subnet-zero** global configuration command to restore the default and disable the use of subnet zero.

Classless Routing

By default, classless routing behavior is enabled on the switch when it is configured to route. With classless routing, if a router receives packets for a subnet of a network with no default route, the router forwards the packet to the best supernet route. A *supernet* consists of contiguous blocks of Class C address spaces used to simulate a single, larger address space and is designed to relieve the pressure on the rapidly depleting Class B address space.

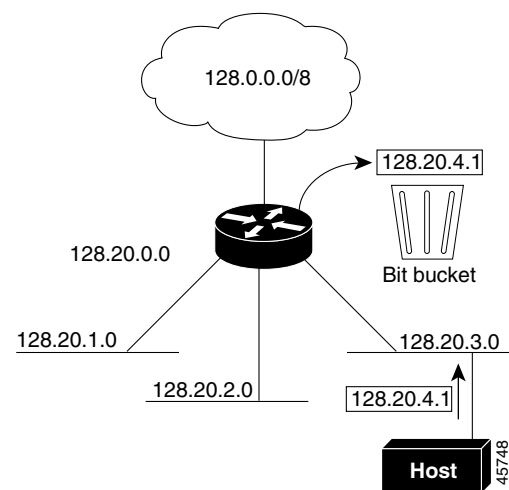
In [Figure 32-2](#), classless routing is enabled. When the host sends a packet to 120.20.4.1, instead of discarding the packet, the router forwards it to the best supernet route. If you disable classless routing and a router receives packets destined for a subnet of a network with no network default route, the router discards the packet.

Figure 32-2 IP Classless Routing



In [Figure 32-3](#), the router in network 128.20.0.0 is connected to subnets 128.20.1.0, 128.20.2.0, and 128.20.3.0. If the host sends a packet to 120.20.4.1, because there is no network default route, the router discards the packet.

Figure 32-3 No IP Classless Routing



To prevent the switch from forwarding packets destined for unrecognized subnets to the best supernet route possible, you can disable classless routing behavior.

Beginning in privileged EXEC mode, follow these steps to disable classless routing:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	no ip classless	Disable classless routing behavior.
Step 3	end	Return to privileged EXEC mode.
Step 4	show running-config	Verify your entry.
Step 5	copy running-config startup-config	(Optional) Save your entry in the configuration file.

To restore the default and have the switch forward packets destined for a subnet of a network with no network default route to the best supernet route possible, use the **ip classless** global configuration command.

Configuring Address Resolution Methods

You can control interface-specific handling of IP by using address resolution. A device using IP can have both a local address or MAC address, which uniquely defines the device on its local segment or LAN, and a network address, which identifies the network to which the device belongs. The local address or MAC address is known as a data link address because it is contained in the data link layer (Layer 2) section of the packet header and is read by data link (Layer 2) devices. To communicate with a device on Ethernet, the software must determine the MAC address of the device. The process of determining the MAC address from an IP address is called *address resolution*. The process of determining the IP address from the MAC address is called *reverse address resolution*.

The switch can use these forms of address resolution:

- Address Resolution Protocol (ARP) is used to associate IP address with MAC addresses. Taking an IP address as input, ARP determines the associated MAC address and then stores the IP address/MAC address association in an ARP cache for rapid retrieval. Then the IP datagram is encapsulated in a link-layer frame and sent over the network. Encapsulation of IP datagrams and ARP requests or replies on IEEE 802 networks other than Ethernet is specified by the Subnetwork Access Protocol (SNAP).
- Proxy ARP helps hosts with no routing tables determine the MAC addresses of hosts on other networks or subnets. If the switch (router) receives an ARP request for a host that is not on the same interface as the ARP request sender, and if the router has all of its routes to the host through other interfaces, it generates a proxy ARP packet giving its own local data link address. The host that sent the ARP request then sends its packets to the router, which forwards them to the intended host.

The switch also uses the Reverse Address Resolution Protocol (RARP), which functions the same as ARP does, except that the RARP packets request an IP address instead of a local MAC address. Using RARP requires a RARP server on the same network segment as the router interface. Use the **ip rarp-server address** interface configuration command to identify the server.

For more information on RARP, see the *Cisco IOS Configuration Fundamentals Configuration Guide, Release 12.2*.

You can perform these tasks to configure address resolution:

- [Define a Static ARP Cache, page 32-9](#)
- [Set ARP Encapsulation, page 32-10](#)
- [Enable Proxy ARP, page 32-10](#)

Define a Static ARP Cache

ARP and other address resolution protocols provide dynamic mapping between IP addresses and MAC addresses. Because most hosts support dynamic address resolution, you usually do not need to specify static ARP cache entries. If you must define a static ARP cache entry, you can do so globally, which installs a permanent entry in the ARP cache that the switch uses to translate IP addresses into MAC addresses. Optionally, you can also specify that the switch respond to ARP requests as if it were the owner of the specified IP address. If you do not want the ARP entry to be permanent, you can specify a timeout period for the ARP entry.

Beginning in privileged EXEC mode, follow these steps to provide static mapping between IP addresses and MAC addresses:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	arp <i>ip-address hardware-address type</i>	Globally associate an IP address with a MAC (hardware) address in the ARP cache, and specify encapsulation type as one of these: • arpa—ARP encapsulation for Ethernet interfaces • snap—Subnetwork Address Protocol encapsulation for Token Ring and FDDI interfaces • sap—HP's ARP type
Step 3	arp <i>ip-address hardware-address type</i> [alias]	(Optional) Specify that the switch respond to ARP requests as if it were the owner of the specified IP address.
Step 4	interface <i>interface-id</i>	Enter interface configuration mode, and specify the interface to configure.
Step 5	arp timeout <i>seconds</i>	(Optional) Set the length of time an ARP cache entry will stay in the cache. The default is 14400 seconds (4 hours). The range is 0 to 2147483 seconds. Note We recommend that you do not set an ARP timeout value lower than 120 seconds.
Step 6	end	Return to privileged EXEC mode.
Step 7	show interfaces [<i>interface-id</i>]	Verify the type of ARP and the timeout value used on all interfaces or a specific interface.
Step 8	show arp show ip arp	View the contents of the ARP cache.
Step 9	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To remove an entry from the ARP cache, use the **no arp** *ip-address hardware-address type* global configuration command. To remove all nonstatic entries from the ARP cache, use the **clear arp-cache** privileged EXEC command.

Set ARP Encapsulation

By default, Ethernet ARP encapsulation (represented by the **arpa** keyword) is enabled on an IP interface. You can change the encapsulation methods to SNAP if required by your network.

Beginning in privileged EXEC mode, follow these steps to specify the ARP encapsulation type:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the Layer 3 interface to configure.
Step 3	arp { arpa snap }	Specify the ARP encapsulation method: • arpa—Address Resolution Protocol • snap—Subnetwork Address Protocol
Step 4	end	Return to privileged EXEC mode.
Step 5	show interfaces [<i>interface-id</i>]	Verify ARP encapsulation configuration on all interfaces or the specified interface.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To disable an encapsulation type, use the **no arp arpa** or **no arp snap** interface configuration command.

Enable Proxy ARP

By default, the switch uses proxy ARP to help hosts determine MAC addresses of hosts on other networks or subnets.

Beginning in privileged EXEC mode, follow these steps to enable proxy ARP if it has been disabled:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the Layer 3 interface to configure.
Step 3	ip proxy-arp	Enable proxy ARP on the interface.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip interface [<i>interface-id</i>]	Verify the configuration on the interface or all interfaces.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To disable proxy ARP on the interface, use the **no ip proxy-arp** interface configuration command.

Routing Assistance When IP Routing is Disabled

These mechanisms allow the switch to learn about routes to other networks when it does not have IP routing enabled:

- [Proxy ARP, page 32-11](#)
- [Default Gateway, page 32-11](#)
- [ICMP Router Discovery Protocol \(IRDP\), page 32-12](#)

Proxy ARP

Proxy ARP, the most common method for learning about other routes, enables an Ethernet host with no routing information to communicate with hosts on other networks or subnets. The host assumes that all hosts are on the same local Ethernet and that they can use ARP to determine their MAC addresses. If a switch receives an ARP request for a host that is not on the same network as the sender, the switch evaluates whether it has the best route to that host. If it does, it sends an ARP reply packet with its own Ethernet MAC address, and the host that sent the request sends the packet to the switch, which forwards it to the intended host. Proxy ARP treats all networks as if they are local and performs ARP requests for every IP address.

Proxy ARP is enabled by default. To enable it after it has been disabled, see the [“Enable Proxy ARP” section on page 32-10](#). Proxy ARP works as long as other routers support it.

Default Gateway

Another method for locating routes is to define a default router or default gateway. All nonlocal packets are sent to this router, which either routes them appropriately or sends an IP Control Message Protocol (ICMP) redirect message back, defining which local router the host should use. The switch caches the redirect messages and forwards each packet as efficiently as possible. A limitation of this method is that there is no means of detecting when the default router has gone down or is unavailable.

Beginning in privileged EXEC mode, follow these steps to define a default gateway (router) when IP routing is disabled:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip default-gateway <i>ip-address</i>	Set up a default gateway (router).
Step 3	end	Return to privileged EXEC mode.
Step 4	show ip redirects	Display the address of the default gateway router to verify the setting.
Step 5	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no ip default-gateway** global configuration command to disable this function.

This example shows how to set and verify a default gateway:

```
Switch(config)# ip default-gateway 10.1.5.59
Switch(config)# end
Switch# show ip redirect
Default gateway is 10.1.5.59
Host          Gateway          Last Use    Total Uses  Interface
ICMP redirect cache is empty
```

ICMP Router Discovery Protocol (IRDP)

Router discovery allows the switch to dynamically learn about routes to other networks using IRDP. IRDP allows hosts to locate routers. When operating as a client, the switch generates router discovery packets. When operating as a host, the switch receives router discovery packets. The switch can also listen to Routing Information Protocol (RIP) routing updates and use this information to infer locations of routers. The switch does not actually store the routing tables sent by routing devices; it merely keeps track of which systems are sending the data. The advantage of using IRDP is that it allows each router to specify both a priority and the time after which a device is assumed to be down if no further packets are received.

Each device discovered becomes a candidate for the default router, and a new highest-priority router is selected when a higher priority router is discovered, when the current default router is declared down, or when a TCP connection is about to time out because of excessive retransmissions.

The only required task for IRDP routing on an interface is to enable IRDP processing on that interface. When enabled, the default parameters apply. You can optionally change any of these parameters.

Beginning in privileged EXEC mode, follow these steps to enable and configure IRDP on an interface:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the Layer 3 interface to configure.
Step 3	ip irdp	Enable IRDP processing on the interface.
Step 4	ip irdp multicast	(Optional) Send IRDP advertisements to the multicast address (224.0.0.1) instead of IP broadcasts. Note This command allows for compatibility with Sun Microsystems Solaris, which requires IRDP packets to be sent out as multicasts. Many implementations cannot receive these multicasts; ensure end-host ability before using this command.
Step 5	ip irdp holdtime <i>seconds</i>	(Optional) Set the IRDP period for which advertisements are valid. The default is three times the maxadvertinterval value. It must be greater than maxadvertinterval and cannot be greater than 9000 seconds. If you change the maxadvertinterval value, this value also changes.
Step 6	ip irdp maxadvertinterval <i>seconds</i>	(Optional) Set the IRDP maximum interval between advertisements. The default is 600 seconds.
Step 7	ip irdp minadvertinterval <i>seconds</i>	(Optional) Set the IRDP minimum interval between advertisements. The default is 0.75 times the maxadvertinterval . If you change the maxadvertinterval , this value changes to the new default (0.75 of maxadvertinterval).
Step 8	ip irdp preference <i>number</i>	(Optional) Set a device IRDP preference level. The allowed range is -2^{31} to 2^{31} . The default is 0. A higher value increases the router preference level.
Step 9	ip irdp address <i>address</i> [<i>number</i>]	(Optional) Specify an IRDP address and preference to proxy-advertise.
Step 10	end	Return to privileged EXEC mode.
Step 11	show ip irdp	Verify settings by displaying IRDP values.
Step 12	copy running-config startup-config	(Optional) Save your entries in the configuration file.

If you change the **maxadvertinterval** value, the **holdtime** and **minadvertinterval** values also change, so it is important to first change the **maxadvertinterval** value, before manually changing either the **holdtime** or **minadvertinterval** values.

Use the **no ip irdp** interface configuration command to disable IRDP routing.

Configuring Broadcast Packet Handling

After configuring an IP interface address, you can enable routing and configure one or more routing protocols, or you can configure the way the switch responds to network broadcasts. A broadcast is a data packet destined for all hosts on a physical network. The switch supports two kinds of broadcasting:

- A directed broadcast packet is sent to a specific network or series of networks. A directed broadcast address includes the network or subnet fields.
- A flooded broadcast packet is sent to every network.



Note

You can also limit broadcast, unicast, and multicast traffic on Layer 2 interfaces by using the **storm-control** interface configuration commands. For more information, see [Chapter 22, “Configuring Port-Based Traffic Control.”](#)

Routers provide some protection from broadcast storms by limiting their extent to the local cable. Bridges (including intelligent bridges), because they are Layer 2 devices, forward broadcasts to all network segments, thus propagating broadcast storms. The best solution to the broadcast storm problem is to use a single broadcast address scheme on a network. In most modern IP implementations, you can set the address to be used as the broadcast address. Many implementations, including the one in the switch, support several addressing schemes for forwarding broadcast messages.

Perform the tasks in these sections to enable these schemes:

- [Enabling Directed Broadcast-to-Physical Broadcast Translation, page 32-13](#)
- [Forwarding UDP Broadcast Packets and Protocols, page 32-14](#)
- [Establishing an IP Broadcast Address, page 32-15](#)
- [Flooding IP Broadcasts, page 32-16](#)

Enabling Directed Broadcast-to-Physical Broadcast Translation

By default, IP directed broadcasts are dropped; they are not forwarded. Dropping IP-directed broadcasts makes routers less susceptible to denial-of-service attacks.

You can enable forwarding of IP-directed broadcasts on an interface where the broadcast becomes a physical (MAC-layer) broadcast. Only those protocols configured by using the **ip forward-protocol** global configuration command are forwarded.

You can specify an access list to control which broadcasts are forwarded. When an access list is specified, only those IP packets permitted by the access list are eligible to be translated from directed broadcasts to physical broadcasts. For more information on access lists, see [Chapter 29, “Configuring Network Security with ACLs.”](#)

Beginning in privileged EXEC mode, follow these steps to enable forwarding of IP-directed broadcasts on an interface:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the interface to configure.
Step 3	ip directed-broadcast [<i>access-list-number</i>]	Enable directed broadcast-to-physical broadcast translation on the interface. You can include an access list to control which broadcasts are forwarded. When you specify an access list, only IP packets permitted by the access list can be translated. Note The ip directed-broadcast interface configuration command can be configured on a VPN routing/forwarding (VRF) interface and is VRF aware. Directed broadcast traffic is routed only within the VRF.
Step 4	exit	Return to global configuration mode.
Step 5	ip forward-protocol { udp [<i>port</i>] nd sdns }	Specify which protocols and ports the router forwards when forwarding broadcast packets. • udp—Forward UDP datagrams. <i>port</i>: (Optional) Destination port that controls which UDP services are forwarded. • nd—Forward ND datagrams. • sdns—Forward SDNS datagrams
Step 6	end	Return to privileged EXEC mode.
Step 7	show ip interface [<i>interface-id</i>] show running-config	Verify the configuration on the interface or all interfaces.
Step 8	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no ip directed-broadcast** interface configuration command to disable translation of directed broadcast to physical broadcasts. Use the **no ip forward-protocol** global configuration command to remove a protocol or port.

Forwarding UDP Broadcast Packets and Protocols

User Datagram Protocol (UDP) is an IP host-to-host layer protocol, as is TCP. UDP provides a low-overhead, connectionless session between two end systems and does not provide for acknowledgment of received datagrams. Network hosts occasionally use UDP broadcasts to determine address, configuration, and name information. If such a host is on a network segment that does not include a server, UDP broadcasts are normally not forwarded. You can remedy this situation by configuring an interface on a router to forward certain classes of broadcasts to a helper address. You can use more than one helper address per interface.

You can specify a UDP destination port to control which UDP services are forwarded. You can specify multiple UDP protocols. You can also specify the Network Disk (ND) protocol, which is used by older diskless Sun workstations and the network security protocol SDNS.

By default, both UDP and ND forwarding are enabled if a helper address has been defined for an interface. The description for the **ip forward-protocol** interface configuration command in the *Cisco IOS IP Command Reference, Volume 1 of 3: Addressing and Services, Release 12.2* lists the ports that are forwarded by default if you do not specify any UDP ports.

If you do not specify any UDP ports when you configure the forwarding of UDP broadcasts, you are configuring the router to act as a BOOTP forwarding agent. BOOTP packets carry DHCP information.

Beginning in privileged EXEC mode, follow these steps to enable forwarding UDP broadcast packets on an interface and specify the destination address:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the Layer 3 interface to configure.
Step 3	ip helper-address <i>address</i>	Enable forwarding and specify the destination address for forwarding UDP broadcast packets, including BOOTP.
Step 4	exit	Return to global configuration mode.
Step 5	ip forward-protocol { udp [<i>port</i>] nd sdns }	Specify which protocols the router forwards when forwarding broadcast packets.
Step 6	end	Return to privileged EXEC mode.
Step 7	show ip interface [<i>interface-id</i>] show running-config	Verify the configuration on the interface or all interfaces.
Step 8	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no ip helper-address** interface configuration command to disable the forwarding of broadcast packets to specific addresses. Use the **no ip forward-protocol** global configuration command to remove a protocol or port.

Establishing an IP Broadcast Address

The most popular IP broadcast address (and the default) is an address consisting of all ones (255.255.255.255). However, the switch can be configured to generate any form of IP broadcast address.

Beginning in privileged EXEC mode, follow these steps to set the IP broadcast address on an interface:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the interface to configure.
Step 3	ip broadcast-address <i>ip-address</i>	Enter a broadcast address different from the default, for example 128.1.255.255.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip interface [<i>interface-id</i>]	Verify the broadcast address on the interface or all interfaces.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To restore the default IP broadcast address, use the **no ip broadcast-address** interface configuration command.

Flooding IP Broadcasts

You can allow IP broadcasts to be flooded throughout your internetwork in a controlled fashion by using the database created by the bridging STP. Using this feature also prevents loops. To support this capability, bridging must be configured on each interface that is to participate in the flooding. If bridging is not configured on an interface, it still can receive broadcasts. However, the interface never forwards broadcasts it receives, and the router never uses that interface to send broadcasts received on a different interface.

Packets that are forwarded to a single network address using the IP helper-address mechanism can be flooded. Only one copy of the packet is sent on each network segment.

To be considered for flooding, packets must meet these criteria. (Note that these are the same conditions used to consider packet forwarding using IP helper addresses.)

- The packet must be a MAC-level broadcast.
- The packet must be an IP-level broadcast.
- The packet must be a TFTP, DNS, Time, NetBIOS, ND, or BOOTP packet, or a UDP specified by the **ip forward-protocol udp** global configuration command.
- The time-to-live (TTL) value of the packet must be at least two.

A flooded UDP datagram is given the destination address specified with the **ip broadcast-address** interface configuration command on the output interface. The destination address can be set to any address. Thus, the destination address might change as the datagram propagates through the network. The source address is never changed. The TTL value is decremented.

When a flooded UDP datagram is sent out an interface (and the destination address possibly changed), the datagram is handed to the normal IP output routines and is, therefore, subject to access lists, if they are present on the output interface.

Beginning in privileged EXEC mode, follow these steps to use the bridging spanning-tree database to flood UDP datagrams:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip forward-protocol spanning-tree	Use the bridging spanning-tree database to flood UDP datagrams.
Step 3	end	Return to privileged EXEC mode.
Step 4	show running-config	Verify your entry.
Step 5	copy running-config startup-config	(Optional) Save your entry in the configuration file.

Use the **no ip forward-protocol spanning-tree** global configuration command to disable the flooding of IP broadcasts.

In the switch, the majority of packets are forwarded in hardware; most packets do not go through the switch CPU. For those packets that do go to the CPU, you can speed up spanning tree-based UDP flooding by a factor of about four to five times by using turbo-flooding. This feature is supported over Ethernet interfaces configured for ARP encapsulation.

Beginning in privileged EXEC mode, follow these steps to increase spanning-tree-based flooding:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode
Step 2	ip forward-protocol turbo-flood	Use the spanning-tree database to speed up flooding of UDP datagrams.
Step 3	end	Return to privileged EXEC mode.
Step 4	show running-config	Verify your entry.
Step 5	copy running-config startup-config	(Optional) Save your entry in the configuration file.

To disable this feature, use the **no ip forward-protocol turbo-flood** global configuration command.

Monitoring and Maintaining IP Addressing

When the contents of a particular cache, table, or database have become or are suspected to be invalid, you can remove all its contents by using the **clear** privileged EXEC commands. [Table 32-3](#) lists the commands for clearing contents.

Table 32-3 Commands to Clear Caches, Tables, and Databases

Command	Purpose
clear arp-cache	Clear the IP ARP cache and the fast-switching cache.
clear host { <i>name</i> *}	Remove one or all entries from the host name and the address cache.
clear ip route { <i>network</i> [<i>mask</i>] *}	Remove one or more routes from the IP routing table.

You can display specific statistics, such as the contents of IP routing tables, caches, and databases; the reachability of nodes; and the routing path that packets are taking through the network. [Table 32-4](#) lists the privileged EXEC commands for displaying IP statistics.

Table 32-4 Commands to Display Caches, Tables, and Databases

Command	Purpose
show arp	Display the entries in the ARP table.
show hosts	Display the default domain name, style of lookup service, name server hosts, and the cached list of host names and addresses.
show ip aliases	Display IP addresses mapped to TCP ports (aliases).
show ip arp	Display the IP ARP cache.
show ip interface [<i>interface-id</i>]	Display the IP status of interfaces.
show ip irdp	Display IRDP values.
show ip masks <i>address</i>	Display the masks used for network addresses and the number of subnets using each mask.
show ip redirects	Display the address of a default gateway.
show ip route [<i>address</i> [<i>mask</i>]] [<i>protocol</i>]	Display the current state of the routing table.
show ip route summary	Display the current state of the routing table in summary form.

Enabling IP Unicast Routing

By default, the switch is in Layer 2 switching mode and IP routing is disabled. To use the Layer 3 capabilities of the switch, you must enable IP routing.

Beginning in privileged EXEC mode, follow these steps to enable IP routing:

	Command	Purpose
Step 1	<code>configure terminal</code>	Enter global configuration mode.
Step 2	<code>ip routing</code>	Enable IP routing.
Step 3	<code>router ip_routing_protocol</code>	Specify an IP routing protocol. This step might include other commands, such as specifying the networks to route with the network (RIP) router configuration command. For information on specific protocols, see sections later in this chapter and in the <i>Cisco IOS IP Configuration Guide, Release 12.2</i> . Note The IP base image supports only RIP as a routing protocol.
Step 4	<code>end</code>	Return to privileged EXEC mode.
Step 5	<code>show running-config</code>	Verify your entries.
Step 6	<code>copy running-config startup-config</code>	(Optional) Save your entries in the configuration file.

Use the **no ip routing** global configuration command to disable routing.

This example shows how to enable IP routing using RIP as the routing protocol:

```
Switch# configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)# ip routing
Switch(config)# router rip
Switch(config-router)# network 10.0.0.0
Switch(config-router)# end
```

You can now set up parameters for the selected routing protocols as described in these sections:

- [Configuring RIP, page 32-19](#)
- [Configuring OSPF, page 32-24](#)
- [Configuring EIGRP, page 32-34](#)
- [Configuring BGP, page 32-41](#)

Configuring RIP

The Routing Information Protocol (RIP) is an interior gateway protocol (IGP) created for use in small, homogeneous networks. It is a distance-vector routing protocol that uses broadcast User Datagram Protocol (UDP) data packets to exchange routing information. The protocol is documented in RFC 1058. You can find detailed information about RIP in *IP Routing Fundamentals*, published by Cisco Press.



Note

RIP is the only routing protocol supported by the IP base image; other routing protocols require the IP services image on the switch.

Using RIP, the switch sends routing information updates (advertisements) every 30 seconds. If a router does not receive an update from another router for 180 seconds or more, it marks the routes served by that router as unusable. If there is still no update after 240 seconds, the router removes all routing table entries for the non-updating router.

RIP uses hop counts to rate the value of different routes. The hop count is the number of routers that can be traversed in a route. A directly connected network has a hop count of zero; a network with a hop count of 16 is unreachable. This small range (0 to 15) makes RIP unsuitable for large networks.

If the router has a default network path, RIP advertises a route that links the router to the pseudonetwork 0.0.0.0. The 0.0.0.0 network does not exist; it is treated by RIP as a network to implement the default routing feature. The switch advertises the default network if a default was learned by RIP or if the router has a gateway of last resort and RIP is configured with a default metric. RIP sends updates to the interfaces in specified networks. If an interface's network is not specified, it is not advertised in any RIP update.

This section briefly describes how to configure RIP. It includes this information:

- [Default RIP Configuration, page 32-19](#)
- [Configuring Basic RIP Parameters, page 32-20](#)
- [Configuring RIP Authentication, page 32-22](#)
- [Configuring Summary Addresses and Split Horizon, page 32-22](#)

Default RIP Configuration

[Table 32-5](#) shows the default RIP configuration.

Table 32-5 **Default RIP Configuration**

Feature	Default Setting
Auto summary	Enabled.
Default-information originate	Disabled.
Default metric	Built-in; automatic metric translations.
IP RIP authentication key-chain	No authentication. Authentication mode: clear text.
IP RIP receive version	According to the version router configuration command.
IP RIP send version	According to the version router configuration command.
IP RIP triggered	According to the version router configuration command.

Table 32-5 **Default RIP Configuration (continued)**

Feature	Default Setting
IP split horizon	Varies with media.
Neighbor	None defined.
Network	None specified.
Offset list	Disabled.
Output delay	0 milliseconds.
Timers basic	• Update: 30 seconds. • Invalid: 180 seconds. • Hold-down: 180 seconds. • Flush: 240 seconds.
Validate-update-source	Enabled.
Version	Receives RIP version 1 and 2 packets; sends version 1 packets.

Configuring Basic RIP Parameters

To configure RIP, you enable RIP routing for a network and optionally configure other parameters.

Beginning in privileged EXEC mode, follow these steps to enable and configure RIP:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip routing	Enable IP routing (required only if IP routing is disabled).
Step 3	router rip	Enable a RIP routing process, and enter router configuration mode.
Step 4	network <i>network number</i>	Associate a network with a RIP routing process. You can specify multiple network commands. RIP routing updates are sent and received through interfaces only on these networks.
Step 5	neighbor <i>ip-address</i>	(Optional) Define a neighboring router with which to exchange routing information. This step allows routing updates from RIP (normally a broadcast protocol) to reach nonbroadcast networks.
Step 6	offset list [<i>access-list number name</i>] { in out } <i>offset [type number]</i>	(Optional) Apply an offset list to routing metrics to increase incoming and outgoing metrics to routes learned through RIP. You can limit the offset list with an access list or an interface.

	Command	Purpose
Step 7	timers basic <i>update invalid holddown flush</i>	(Optional) Adjust routing protocol timers. Valid ranges for all timers are 0 to 4294967295 seconds. <i>update</i>—Time between sending routing updates. The default is 30 seconds. <i>invalid</i>—Time after which a route is declared invalid. The default is 180 seconds. <i>holddown</i>—Time before a route is removed from the routing table. The default is 180 seconds. <i>flush</i>—Amount of time for which routing updates are postponed. The default is 240 seconds.
Step 8	version {1 2}	(Optional) Configure the switch to receive and send only RIP Version 1 or RIP version 2 packets. By default, the switch receives Version 1 and 2 but sends only Version 1. You can also use the interface commands ip rip {send receive} version 1 2 1 2 to control what versions are used for sending and receiving on interfaces.
Step 9	no auto summary	(Optional) Disable automatic summarization. By default, the switch summarizes subprefixes when crossing classful network boundaries. Disable summarization (RIP version 2 only) to advertise subnet and host routing information to classful network boundaries.
Step 10	no validate-update-source	(Optional) Disable validation of the source IP address of incoming RIP routing updates. By default, the switch validates the source IP address of incoming RIP routing updates and discards the update if the source address is not valid. Under normal circumstances, disabling this feature is not recommended. However, if you have a router that is off-network and you want to receive its updates, you can use this command.
Step 11	output-delay <i>delay</i>	(Optional) Add interpacket delay for RIP updates sent. By default, packets in a multiple-packet RIP update have no delay added between packets. If you are sending packets to a lower-speed device, you can add an interpacket delay in the range of 8 to 50 milliseconds.
Step 12	end	Return to privileged EXEC mode.
Step 13	show ip protocols	Verify your entries.
Step 14	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To turn off the RIP routing process, use the **no router rip** global configuration command.

To display the parameters and current state of the active routing protocol process, use the **show ip protocols** privileged EXEC command. Use the **show ip rip database** privileged EXEC command to display summary address entries in the RIP database.

Configuring RIP Authentication

RIP version 1 does not support authentication. If you are sending and receiving RIP Version 2 packets, you can enable RIP authentication on an interface. The key chain determines the set of keys that can be used on the interface. If a key chain is not configured, no authentication is performed, not even the default. Therefore, you must also perform the tasks in the “[Managing Authentication Keys](#)” section on page 32-84.

The switch supports two modes of authentication on interfaces for which RIP authentication is enabled: plain text and MD5. The default is plain text.

Beginning in privileged EXEC mode, follow these steps to configure RIP authentication on an interface:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the interface to configure.
Step 3	ip rip authentication key-chain <i>name-of-chain</i>	Enable RIP authentication.
Step 4	ip rip authentication mode {text md5}	Configure the interface to use plain text authentication (the default) or MD5 digest authentication.
Step 5	end	Return to privileged EXEC mode.
Step 6	show running-config interface [<i>interface-id</i>]	Verify your entries.
Step 7	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To restore clear text authentication, use the **no ip rip authentication mode** interface configuration command. To prevent authentication, use the **no ip rip authentication key-chain** interface configuration command.

Configuring Summary Addresses and Split Horizon

Routers connected to broadcast-type IP networks and using distance-vector routing protocols normally use the split-horizon mechanism to reduce the possibility of routing loops. Split horizon blocks information about routes from being advertised by a router on any interface from which that information originated. This feature usually optimizes communication among multiple routers, especially when links are broken.



Note

In general, disabling split horizon is not recommended unless you are certain that your application requires it to properly advertise routes.

If you want to configure an interface running RIP to advertise a summarized local IP address pool on a network access server for dial-up clients, use the **ip summary-address rip** interface configuration command.



Note

If split horizon is enabled, neither autosummary nor interface IP summary addresses are advertised.

Beginning in privileged EXEC mode, follow these steps to set an interface to advertise a summarized local IP address and to disable split horizon on the interface:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the Layer 3 interface to configure.
Step 3	ip address <i>ip-address subnet-mask</i>	Configure the IP address and IP subnet.
Step 4	ip summary-address rip <i>ip address ip-network mask</i>	Configure the IP address to be summarized and the IP network mask.
Step 5	no ip split horizon	Disable split horizon on the interface.
Step 6	end	Return to privileged EXEC mode.
Step 7	show ip interface <i>interface-id</i>	Verify your entries.
Step 8	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To disable IP summarization, use the **no ip summary-address rip** router configuration command.

In this example, the major net is 10.0.0.0. The summary address 10.2.0.0 overrides the autosummary address of 10.0.0.0 so that 10.2.0.0 is advertised out interface Gigabit Ethernet 0/2, and 10.0.0.0 is not advertised. In the example, if the interface is still in Layer 2 mode (the default), you must enter a **no switchport** interface configuration command before entering the **ip address** interface configuration command.



Note

If split horizon is enabled, neither autosummary nor interface summary addresses (those configured with the **ip summary-address rip** router configuration command) are advertised.

```
Switch(config)# router rip
Switch(config-router)# interface gi0/2
Switch(config-if)# ip address 10.1.5.1 255.255.255.0
Switch(config-if)# ip summary-address rip 10.2.0.0 255.255.0.0
Switch(config-if)# no ip split-horizon
Switch(config-if)# exit
Switch(config)# router rip
Switch(config-router)# network 10.0.0.0
Switch(config-router)# neighbor 2.2.2.2 peer-group mygroup
Switch(config-router)# end
```

Configuring OSPF

This section briefly describes how to configure Open Shortest Path First (OSPF). For a complete description of the OSPF commands, see the “OSPF Commands” chapter of the *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*.

**Note**

OSPF classifies different media into broadcast, nonbroadcast, and point-to-point networks. The switch supports broadcast (Ethernet, Token Ring, and FDDI) and point-to-point networks (Ethernet interfaces configured as point-to-point links).

OSPF is an Interior Gateway Protocol (IGP) designed expressly for IP networks, supporting IP subnetting and tagging of externally derived routing information. OSPF also allows packet authentication and uses IP multicast when sending and receiving packets. The Cisco implementation supports RFC 1253, OSPF management information base (MIB).

The Cisco implementation conforms to the OSPF Version 2 specifications with these key features:

- Definition of stub areas is supported.
- Routes learned through any IP routing protocol can be redistributed into another IP routing protocol. At the intradomain level, this means that OSPF can import routes learned through RIP. OSPF routes can also be exported into RIP.
- Plain text and MD5 authentication among neighboring routers within an area is supported.
- Configurable routing interface parameters include interface output cost, retransmission interval, interface transmit delay, router priority, router dead and hello intervals, and authentication key.
- Virtual links are supported.
- Not-so-stubby-areas (NSSAs) per RFC 1587 are supported.

OSPF typically requires coordination among many internal routers, *area border routers* (ABRs) connected to multiple areas, and *autonomous system boundary routers* (ASBRs). The minimum configuration would use all default parameter values, no authentication, and interfaces assigned to areas. If you customize your environment, you must ensure coordinated configuration of all routers.

This section briefly describes how to configure OSPF. It includes this information:

- [Default OSPF Configuration, page 32-25](#)
- [Configuring Basic OSPF Parameters, page 32-26](#)
- [Configuring OSPF Interfaces, page 32-27](#)
- [Configuring OSPF Area Parameters, page 32-28](#)
- [Configuring Other OSPF Parameters, page 32-30](#)
- [Changing LSA Group Pacing, page 32-32](#)
- [Configuring a Loopback Interface, page 32-32](#)
- [Monitoring OSPF, page 32-33](#)

Default OSPF Configuration

Table 32-6 shows the default OSPF configuration.

Table 32-6 *Default OSPF Configuration*

Feature	Default Setting
Interface parameters	Cost: No default cost predefined. Retransmit interval: 5 seconds. Transmit delay: 1 second. Priority: 1. Hello interval: 10 seconds. Dead interval: 4 times the hello interval. No authentication. No password specified. MD5 authentication disabled.
Area	Authentication type: 0 (no authentication). Default cost: 1. Range: Disabled. Stub: No stub area defined. NSSA: No NSSA area defined.
Auto cost	100 Mbps.
Default-information originate	Disabled. When enabled, the default metric setting is 10, and the external route type default is Type 2.
Default metric	Built-in, automatic metric translation, as appropriate for each routing protocol.
Distance OSPF	dist1 (all routes within an area): 110. dist2 (all routes from one area to another): 110. and dist3 (routes from other routing domains): 110.
OSPF database filter	Disabled. All outgoing link-state advertisements (LSAs) are flooded to the interface.
IP OSPF name lookup	Disabled.
Log adjacency changes	Enabled.
Neighbor	None specified.
Neighbor database filter	Disabled. All outgoing LSAs are flooded to the neighbor.
Network area	Disabled.
NSF ¹ awareness	Enabled ² . Allows Layer 3 switches to continue forwarding packets from a neighboring NSF-capable router during hardware or software changes.
Router ID	No OSPF routing process defined.
Summary address	Disabled.

Table 32-6 Default OSPF Configuration (continued)

Feature	Default Setting
Timers LSA group pacing	240 seconds.
Timers shortest path first (spf)	spf delay: 5 seconds. spf-holdtime: 10 seconds.
Virtual link	No area ID or router ID defined. Hello interval: 10 seconds. Retransmit interval: 5 seconds. Transmit delay: 1 second. Dead interval: 40 seconds. Authentication key: no key predefined. Message-digest key (MD5): no key predefined.

1. NSF = Nonstop forwarding
2. OSPF NSF awareness is enabled on Catalyst 3550, 3560 and 3750 switches running the IP services image, Cisco IOS Release 12.2(25)SEC or later.

Nonstop Forwarding Awareness

The OSPF NSF awareness feature is supported in the IP services image, beginning with Cisco IOS Release 12.2(25)SEC. When the neighboring router is NSF-capable, the Layer 3 switch continues to forward packets from the neighboring router during the interval between the primary Route Processor (RP) in a router failing and the backup RP taking over, or while the primary RP is manually reloaded for a nondisruptive software upgrade.

This feature cannot be disabled. For more information on this feature, see *OSPF Nonstop Forwarding (NSF) Awareness* at this URL:

http://www.cisco.com/en/US/products/ps6350/products_configuration_guide_chapter09186a00804557a8.html

Configuring Basic OSPF Parameters

Enabling OSPF requires that you create an OSPF routing process, specify the range of IP addresses to be associated with the routing process, and assign area IDs to be associated with that range.

Beginning in privileged EXEC mode, follow these steps to enable OSPF:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip routing	Enable IP routing (required only if IP routing is disabled).
Step 3	router ospf process-id	Enable OSPF routing, and enter router configuration mode. The process ID is an internally used identification parameter that is locally assigned and can be any positive integer. Each OSPF routing process has a unique value.

	Command	Purpose
Step 4	network <i>address wildcard-mask area area-id</i>	Define an interface on which OSPF runs and the area ID for that interface. You can use the wildcard-mask to use a single command to define one or more multiple interfaces to be associated with a specific OSPF area. The area ID can be a decimal value or an IP address.
Step 5	end	Return to privileged EXEC mode.
Step 6	show ip protocols	Verify your entries.
Step 7	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To terminate an OSPF routing process, use the **no router ospf process-id** global configuration command.

This example shows how to configure an OSPF routing process and assign it a process number of 109:

```
Switch(config)# router ospf 109
Switch(config-router)# network 131.108.0.0 255.255.255.0 area 24
Switch(config-router)# end
```

Configuring OSPF Interfaces

You can use the **ip ospf** interface configuration commands to modify interface-specific OSPF parameters. You are not required to modify any of these parameters, but some interface parameters (hello interval, dead interval, and authentication key) must be consistent across all routers in an attached network. If you modify these parameters, be sure all routers in the network have compatible values.



Note

The **ip ospf** interface configuration commands are all optional.

Beginning in privileged EXEC mode, follow these steps to modify OSPF interface parameters:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the Layer 3 interface to configure.
Step 3	ip ospf cost	(Optional) Explicitly specify the cost of sending a packet on the interface.
Step 4	ip ospf retransmit-interval <i>seconds</i>	(Optional) Specify the number of seconds between link state advertisement transmissions. The range is 1 to 65535 seconds. The default is 5 seconds.
Step 5	ip ospf transmit-delay <i>seconds</i>	(Optional) Set the estimated number of seconds to wait before sending a link state update packet. The range is 1 to 65535 seconds. The default is 1 second.
Step 6	ip ospf priority <i>number</i>	(Optional) Set priority to help determine the OSPF designated router for a network. The range is from 0 to 255. The default is 1.
Step 7	ip ospf hello-interval <i>seconds</i>	(Optional) Set the number of seconds between hello packets sent on an OSPF interface. The value must be the same for all nodes on a network. The range is 1 to 65535 seconds. The default is 10 seconds.

	Command	Purpose
Step 8	ip ospf dead-interval <i>seconds</i>	(Optional) Set the number of seconds after the last device hello packet was seen before its neighbors declare the OSPF router to be down. The value must be the same for all nodes on a network. The range is 1 to 65535 seconds. The default is 4 times the hello interval.
Step 9	ip ospf authentication-key <i>key</i>	(Optional) Assign a password to be used by neighboring OSPF routers. The password can be any string of keyboard-entered characters up to 8 bytes in length. All neighboring routers on the same network must have the same password to exchange OSPF information.
Step 10	ip ospf message digest-key <i>keyid md5 key</i>	(Optional) Enable MDS authentication. <i>keyid</i>—An identifier from 1 to 255. <i>key</i>—An alphanumeric password of up to 16 bytes.
Step 11	ip ospf database-filter all out	(Optional) Block flooding of OSPF LSA packets to the interface. By default, OSPF floods new LSAs over all interfaces in the same area, except the interface on which the LSA arrives.
Step 12	end	Return to privileged EXEC mode.
Step 13	show ip ospf interface [<i>interface-name</i>]	Display OSPF-related interface information.
Step 14	show ip ospf neighbor detail	Display NSF awareness status of neighbor switch. The output will match one of the following two examples: <i>Options is 0x52</i> <i>LLS Options is 0x1 (LR)</i> When both of these lines appear, the neighbor switch is NSF aware. <i>Options is 0x42</i>—This means the neighbor switch is not NSF aware.
Step 15	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no** form of these commands to remove the configured parameter value or return to the default value.

Configuring OSPF Area Parameters

You can optionally configure several OSPF area parameters. These parameters include authentication for password-based protection against unauthorized access to an area, stub areas, and not-so-stubby-areas (NSSAs). *Stub areas* are areas into which information on external routes is not sent. Instead, the area border router (ABR) generates a default external route into the stub area for destinations outside the autonomous system (AS). An NSSA does not flood all LSAs from the core into the area, but can import AS external routes within the area by redistribution.

Route summarization is the consolidation of advertised addresses into a single summary route to be advertised by other areas. If network numbers are contiguous, you can use the **area range** router configuration command to configure the ABR to advertise a summary route that covers all networks in the range.

**Note**

The OSPF **area** router configuration commands are all optional.

Beginning in privileged EXEC mode, follow these steps to configure area parameters:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router ospf <i>process-id</i>	Enable OSPF routing, and enter router configuration mode.
Step 3	area <i>area-id</i> authentication	(Optional) Allow password-based protection against unauthorized access to the identified area. The identifier can be either a decimal value or an IP address.
Step 4	area <i>area-id</i> authentication message-digest	(Optional) Enable MD5 authentication on the area.
Step 5	area <i>area-id</i> stub [no-summary]	(Optional) Define an area as a stub area. The no-summary keyword prevents an ABR from sending summary link advertisements into the stub area.
Step 6	area <i>area-id</i> nssa [no-redistribution] [default-information-originate] [no-summary]	(Optional) Defines an area as a not-so-stubby-area. Every router within the same area must agree that the area is NSSA. Select one of these keywords: • no-redistribution—Select when the router is an NSSA ABR and you want the redistribute command to import routes into normal areas, but not into the NSSA. • default-information-originate—Select on an ABR to allow importing type 7 LSAs into the NSSA. • no-redistribution—Select to not send summary LSAs into the NSSA.
Step 7	area <i>area-id</i> range <i>address mask</i>	(Optional) Specify an address range for which a single route is advertised. Use this command only with area border routers.
Step 8	end	Return to privileged EXEC mode.
Step 9	show ip ospf [<i>process-id</i>] show ip ospf [<i>process-id</i> [<i>area-id</i>]] database	Display information about the OSPF routing process in general or for a specific process ID to verify configuration. Display lists of information related to the OSPF database for a specific router.
Step 10	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no** form of these commands to remove the configured parameter value or to return to the default value.

Configuring Other OSPF Parameters

You can optionally configure other OSPF parameters in router configuration mode.

- **Route summarization:** When redistributing routes from other protocols as described in the [“Using Route Maps to Redistribute Routing Information” section on page 32-76](#), each route is advertised individually in an external LSA. To help decrease the size of the OSPF link state database, you can use the **summary-address** router configuration command to advertise a single router for all the redistributed routes included in a specified network address and mask.
- **Virtual links:** In OSPF, all areas must be connected to a backbone area. You can establish a virtual link in case of a backbone-continuity break by configuring two Area Border Routers as endpoints of a virtual link. Configuration information includes the identity of the other virtual endpoint (the other ABR) and the nonbackbone link that the two routers have in common (the transit area). Virtual links cannot be configured through a stub area.
- **Default route:** When you specifically configure redistribution of routes into an OSPF routing domain, the route automatically becomes an autonomous system boundary router (ASBR). You can force the ASBR to generate a default route into the OSPF routing domain.
- **Domain Name Server (DNS) names for use in all OSPF **show** privileged EXEC command displays** makes it easier to identify a router than displaying it by router ID or neighbor ID.
- **Default Metrics:** OSPF calculates the OSPF metric for an interface according to the bandwidth of the interface. The metric is calculated as *ref-bw* divided by bandwidth, where *ref* is 10 by default, and bandwidth (*bw*) is determined by the **bandwidth** interface configuration command. For multiple links with high bandwidth, you can specify a larger number to differentiate the cost on those links.
- **Administrative distance** is a rating of the trustworthiness of a routing information source, an integer between 0 and 255, with a higher value meaning a lower trust rating. An administrative distance of 255 means the routing information source cannot be trusted at all and should be ignored. OSPF uses three different administrative distances: routes within an area (interarea), routes to another area (interarea), and routes from another routing domain learned through redistribution (external). You can change any of the distance values.
- **Passive interfaces:** Because interfaces between two devices on an Ethernet represent only one network segment, to prevent OSPF from sending hello packets for the sending interface, you must configure the sending device to be a passive interface. Both devices can identify each other through the hello packet for the receiving interface.
- **Route calculation timers:** You can configure the delay time between when OSPF receives a topology change and when it starts the shortest path first (SPF) calculation and the hold time between two SPF calculations.
- **Log neighbor changes:** You can configure the router to send a syslog message when an OSPF neighbor state changes, providing a high-level view of changes in the router.

Beginning in privileged EXEC mode, follow these steps to configure these OSPF parameters:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router ospf <i>process-id</i>	Enable OSPF routing, and enter router configuration mode.
Step 3	summary-address <i>address mask</i>	(Optional) Specify an address and IP subnet mask for redistributed routes so that only one summary route is advertised.

	Command	Purpose
Step 4	area <i>area-id</i> virtual-link <i>router-id</i> [hello-interval <i>seconds</i>] [retransmit-interval <i>seconds</i>] [trans] [[authentication-key <i>key</i>] message-digest-key <i>keyid</i> md5 <i>key</i>]]	(Optional) Establish a virtual link and set its parameters. See the “Configuring OSPF Interfaces” section on page 32-27 for parameter definitions and Table 32-6 on page 32-25 for virtual link defaults.
Step 5	default-information originate [always] [metric <i>metric-value</i>] [metric-type <i>type-value</i>] [route-map <i>map-name</i>]	(Optional) Force the ASBR to generate a default route into the OSPF routing domain. Parameters are all optional.
Step 6	ip ospf name-lookup	(Optional) Configure DNS name lookup. The default is disabled.
Step 7	ip auto-cost reference-bandwidth <i>ref-bw</i>	(Optional) Specify an address range for which a single route will be advertised. Use this command only with area border routers.
Step 8	distance ospf {[inter-area <i>dist1</i>] [inter-area <i>dist2</i>] [external <i>dist3</i>]}	(Optional) Change the OSPF distance values. The default distance for each type of route is 110. The range is 1 to 255.
Step 9	passive-interface <i>type number</i>	(Optional) Suppress the sending of hello packets through the specified interface.
Step 10	timers spf <i>spf-delay</i> <i>spf-holdtime</i>	(Optional) Configure route calculation timers. <i>spf-delay</i>—Enter an integer from 0 to 65535. The default is 5 seconds; 0 means no delay. <i>spf-holdtime</i>—Enter an integer from 0 to 65535. The default is 10 seconds; 0 means no delay.
Step 11	ospf log-adj-changes	(Optional) Send syslog message when a neighbor state changes.
Step 12	end	Return to privileged EXEC mode.
Step 13	show ip ospf [<i>process-id</i> [<i>area-id</i>]] database	Display lists of information related to the OSPF database for a specific router. For some of the keyword options, see to the “Monitoring OSPF” section on page 32-33.
Step 14	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Changing LSA Group Pacing

The OSPF LSA group pacing feature allows the router to group OSPF LSAs and pace the refreshing, check-summing, and aging functions for more efficient router use. This feature is enabled by default with a 4-minute default pacing interval, and you will not usually need to modify this parameter. The optimum group pacing interval is inversely proportional to the number of LSAs the router is refreshing, check-summing, and aging. For example, if you have approximately 10,000 LSAs in the database, decreasing the pacing interval would benefit you. If you have a very small database (40 to 100 LSAs), increasing the pacing interval to 10 to 20 minutes might benefit you slightly.

Beginning in privileged EXEC mode, follow these steps to configure OSPF LSA pacing:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router ospf <i>process-id</i>	Enable OSPF routing, and enter router configuration mode.
Step 3	timers lsa-group-pacing <i>seconds</i>	Change the group pacing of LSAs.
Step 4	end	Return to privileged EXEC mode.
Step 5	show running-config	Verify your entries.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To return to the default value, use the **no timers lsa-group-pacing** router configuration command.

Configuring a Loopback Interface

OSPF uses the highest IP address configured on the interfaces as its router ID. If this interface is down or removed, the OSPF process must recalculate a new router ID and resend all its routing information out its interfaces. If a loopback interface is configured with an IP address, OSPF uses this IP address as its router ID, even if other interfaces have higher IP addresses. Because loopback interfaces never fail, this provides greater stability. OSPF automatically prefers a loopback interface over other interfaces, and it chooses the highest IP address among all loopback interfaces.

Beginning in privileged EXEC mode, follow these steps to configure a loopback interface:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface loopback 0	Create a loopback interface, and enter interface configuration mode.
Step 3	ip address <i>address mask</i>	Assign an IP address to this interface.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip interface	Verify your entries.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no interface loopback 0** global configuration command to disable the loopback interface.

Monitoring OSPF

You can display specific statistics such as the contents of IP routing tables, caches, and databases.

Table 32-7 lists some of the privileged EXEC commands for displaying statistics. For more **show ip ospf database** privileged EXEC command options and for explanations of fields in the resulting display, see the *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*.

Table 32-7 Show IP OSPF Statistics Commands

Command	Purpose
show ip ospf [process-id]	Display general information about OSPF routing processes.
show ip ospf [process-id] database [router] [link-state-id] show ip ospf [process-id] database [router] [self-originate] show ip ospf [process-id] database [router] [adv-router [ip-address]] show ip ospf [process-id] database [network] [link-state-id] show ip ospf [process-id] database [summary] [link-state-id] show ip ospf [process-id] database [asbr-summary] [link-state-id] show ip ospf [process-id] database [external] [link-state-id] show ip ospf [process-id area-id] database [database-summary]	Display lists of information related to the OSPF database.
show ip ospf border-routes	Display the internal OSPF routing ABR and ASBR table entries.
show ip ospf interface [interface-name]	Display OSPF-related interface information.
show ip ospf neighbor [interface-name] [neighbor-id] detail	Display OSPF interface neighbor information.
show ip ospf virtual-links	Display OSPF-related virtual links information.

Configuring EIGRP

Enhanced IGRP (EIGRP) is a Cisco proprietary enhanced version of the IGRP. EIGRP uses the same distance vector algorithm and distance information as IGRP; however, the convergence properties and the operating efficiency of EIGRP are significantly improved.

The convergence technology employs an algorithm referred to as the Diffusing Update Algorithm (DUAL), which guarantees loop-free operation at every instant throughout a route computation and allows all devices involved in a topology change to synchronize at the same time. Routers that are not affected by topology changes are not involved in recomputations.

IP EIGRP provides increased network width. With RIP, the largest possible width of your network is 15 hops. Because the EIGRP metric is large enough to support thousands of hops, the only barrier to expanding the network is the transport-layer hop counter. EIGRP increments the transport control field only when an IP packet has traversed 15 routers and the next hop to the destination was learned through EIGRP. When a RIP route is used as the next hop to the destination, the transport control field is incremented as usual.

EIGRP offers these features:

- Fast convergence.
- Incremental updates when the state of a destination changes, instead of sending the entire contents of the routing table, minimizing the bandwidth required for EIGRP packets.
- Protocol-independent neighbor discovery mechanism to learn about neighboring routers.
- Variable-length subnet masks (VLSMs).
- Arbitrary route summarization.
- EIGRP scales to large networks.

EIGRP has these four basic components:

- *Neighbor discovery and recovery* is the process that routers use to dynamically learn of other routers on their directly attached networks. Routers must also discover when their neighbors become unreachable or inoperative. Neighbor discovery and recovery is achieved with low overhead by periodically sending small hello packets. As long as hello packets are received, the Cisco IOS software can determine that a neighbor is alive and functioning. When this status is determined, the neighboring routers can exchange routing information.
- *The reliable transport protocol* is responsible for guaranteed, ordered delivery of EIGRP packets to all neighbors. It supports intermixed transmission of multicast and unicast packets. Some EIGRP packets must be sent reliably, and others need not be. For efficiency, reliability is provided only when necessary. For example, on a multiaccess network that has multicast capabilities (such as Ethernet), it is not necessary to send hellos reliably to all neighbors individually. Therefore, EIGRP sends a single multicast hello with an indication in the packet informing the receivers that the packet need not be acknowledged. Other types of packets (such as updates) require acknowledgment, which is shown in the packet. The reliable transport has a provision to send multicast packets quickly when there are unacknowledged packets pending. Doing so helps ensure that convergence time remains low in the presence of varying speed links.

- The *DUAL finite state machine* embodies the decision process for all route computations. It tracks all routes advertised by all neighbors. DUAL uses the distance information (known as a metric) to select efficient, loop-free paths. DUAL selects routes to be inserted into a routing table based on feasible successors. A successor is a neighboring router used for packet forwarding that has a least-cost path to a destination that is guaranteed not to be part of a routing loop. When there are no feasible successors, but there are neighbors advertising the destination, a recomputation must occur. This is the process whereby a new successor is determined. The amount of time it takes to recompute the route affects the convergence time. Recomputation is processor-intensive; it is advantageous to avoid recomputation if it is not necessary. When a topology change occurs, DUAL tests for feasible successors. If there are feasible successors, it uses any it finds to avoid unnecessary recomputation.
- The *protocol-dependent modules* are responsible for network layer protocol-specific tasks. An example is the IP EIGRP module, which is responsible for sending and receiving EIGRP packets that are encapsulated in IP. It is also responsible for parsing EIGRP packets and informing DUAL of the new information received. EIGRP asks DUAL to make routing decisions, but the results are stored in the IP routing table. EIGRP is also responsible for redistributing routes learned by other IP routing protocols.

This section briefly describes how to configure EIGRP. It includes this information:

- [Default EIGRP Configuration, page 32-35](#)
- [Configuring Basic EIGRP Parameters, page 32-37](#)
- [Configuring EIGRP Interfaces, page 32-38](#)
- [Configuring EIGRP Route Authentication, page 32-39](#)
- [Monitoring and Maintaining EIGRP, page 32-40](#)

Default EIGRP Configuration

[Table 32-8](#) shows the default EIGRP configuration.

Table 32-8 **Default EIGRP Configuration**

Feature	Default Setting
Auto summary	Enabled. Subprefixes are summarized to the classful network boundary when crossing classful network boundaries.
Default-information	Exterior routes are accepted and default information is passed between EIGRP processes when doing redistribution.
Default metric	Only connected routes and interface static routes can be redistributed without a default metric. The metric includes: • Bandwidth: 0 or greater kbps. • Delay (tens of microseconds): 0 or any positive number that is a multiple of 39.1 nanoseconds. • Reliability: any number between 0 and 255 (255 means 100 percent reliability). • Loading: effective bandwidth as a number between 0 and 255 (255 is 100 percent loading). • MTU: maximum transmission unit size of the route in bytes. 0 or any positive integer.

Table 32-8 **Default EIGRP Configuration (continued)**

Feature	Default Setting
Distance	Internal distance: 90. External distance: 170.
EIGRP log-neighbor changes	Disabled. No adjacency changes logged.
IP authentication key-chain	No authentication provided.
IP authentication mode	No authentication provided.
IP bandwidth-percent	50 percent.
IP hello interval	For low-speed nonbroadcast multiaccess (NBMA) networks: 60 seconds; all other networks: 5 seconds.
IP hold-time	For low-speed NBMA networks: 180 seconds; all other networks: 15 seconds.
IP split-horizon	Enabled.
IP summary address	No summary aggregate addresses are predefined.
Metric weights	tos: 0; k1 and k3: 1; k2, k4, and k5: 0
Network	None specified.
NSF ¹ Awareness	Enabled ² . Allows Layer 3 switches to continue forwarding packets from a neighboring NSF-capable router during hardware or software changes.
Offset-list	Disabled.
Router EIGRP	Disabled.
Set metric	No metric set in the route map.
Traffic-share	Distributed proportionately to the ratios of the metrics.
Variance	1 (equal-cost load balancing).

1. NSF = Nonstop Forwarding

2. EIGRP NSF awareness is enabled on Catalyst 3550, 3560 and 3750 switches running the IP services image, Cisco IOS Release 12.2(25)SEC or later.

To create an EIGRP routing process, you must enable EIGRP and associate networks. EIGRP sends updates to the interfaces in the specified networks. If you do not specify an interface network, it is not advertised in any EIGRP update.

**Note**

If you have routers on your network that are configured for IGRP, and you want to change to EIGRP, you must designate transition routers that have both IGRP and EIGRP configured. In these cases, perform Steps 1 through 3 in the next section. You must use the same AS number for routes to be automatically redistributed.

Nonstop Forwarding Awareness

The EIGRP NSF Awareness feature is supported in the IP services image, beginning with Cisco IOS Release 12.2(25)SEC. When the neighboring router is NSF-capable, the Layer 3 switch continues to forward packets from the neighboring router during the interval between the primary Route Processor (RP) in a router crashing and the backup RP taking over, or while the primary RP is manually reloaded for a non-disruptive software upgrade.

This feature cannot be disabled. For more information on this feature see *EIGRP Nonstop Forwarding (NSF) Awareness* at this URL:

http://www.cisco.com/en/US/products/sw/iosswrel/ps1839/products_feature_guide09186a0080160010.html

Configuring Basic EIGRP Parameters

Beginning in privileged EXEC mode, follow these steps to configure EIGRP. Configuring the routing process is required; other steps are optional:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip routing	Enable IP routing (required only if IP routing is disabled).
Step 3	router eigrp <i>autonomous-system</i>	Enable an EIGRP routing process, and enter router configuration mode. The AS number identifies the routes to other EIGRP routers and is used to tag routing information.
Step 4	network <i>network-number</i>	Associate networks with an EIGRP routing process. EIGRP sends updates to the interfaces in the specified networks. If an interface's network is not specified, it is not advertised in any EIGRP update.
Step 5	eigrp log-neighbor-changes	(Optional) Enable logging of EIGRP neighbor changes to monitor routing system stability.
Step 6	metric weights <i>tos k1 k2 k3 k4 k5</i>	(Optional) Adjust the EIGRP metric. Although the defaults have been carefully determined to provide excellent operation in most networks, you can adjust them.  Caution Determining metrics is complex and is not recommended without guidance from an experienced network designer.
Step 7	offset list [<i>access-list number name</i>] { in out } <i>offset [type number]</i>	(Optional) Apply an offset list to routing metrics to increase incoming and outgoing metrics to routes learned through EIGRP. You can limit the offset list with an access list or an interface.
Step 8	no auto-summary	(Optional) Disable automatic summarization of subnet routes into network-level routes.
Step 9	ip summary-address eigrp <i>autonomous-system-number address mask</i>	(Optional) Configure a summary aggregate.
Step 10	end	Return to privileged EXEC mode.

	Command	Purpose
Step 11	show ip protocols	Verify your entries. For NSF awareness, the output shows: *** IP Routing is NSF aware *** EIGRP NSF enabled
Step 12	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no** forms of these commands to disable the feature or return the setting to the default value.

Configuring EIGRP Interfaces

Other optional EIGRP parameters can be configured on an interface basis.

Beginning in privileged EXEC mode, follow these steps to configure EIGRP interfaces:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the Layer 3 interface to configure.
Step 3	ip bandwidth-percent eigrp <i>percent</i>	(Optional) Configure the percentage of bandwidth that can be used by EIGRP on an interface. The default is 50 percent.
Step 4	ip summary-address eigrp <i>autonomous-system-number address mask</i>	(Optional) Configure a summary aggregate address for a specified interface (not usually necessary if auto-summary is enabled).
Step 5	ip hello-interval eigrp <i>autonomous-system-number seconds</i>	(Optional) Change the hello time interval for an EIGRP routing process. The range is 1 to 65535 seconds. The default is 60 seconds for low-speed NBMA networks and 5 seconds for all other networks.
Step 6	ip hold-time eigrp <i>autonomous-system-number seconds</i>	(Optional) Change the hold time interval for an EIGRP routing process. The range is 1 to 65535 seconds. The default is 180 seconds for low-speed NBMA networks and 15 seconds for all other networks.  Caution Do not adjust the hold time without consulting Cisco technical support.
Step 7	no ip split-horizon eigrp <i>autonomous-system-number</i>	(Optional) Disable split horizon to allow route information to be advertised by a router out any interface from which that information originated.
Step 8	end	Return to privileged EXEC mode.
Step 9	show ip eigrp interface	Display which interfaces EIGRP is active on and information about EIGRP relating to those interfaces.
Step 10	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no** forms of these commands to disable the feature or return the setting to the default value.

Configuring EIGRP Route Authentication

EIGRP route authentication provides MD5 authentication of routing updates from the EIGRP routing protocol to prevent the introduction of unauthorized or false routing messages from unapproved sources.

Beginning in privileged EXEC mode, follow these steps to enable authentication:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the Layer 3 interface to configure.
Step 3	ip authentication mode eigrp <i>autonomous-system</i> md5	Enable MD5 authentication in IP EIGRP packets.
Step 4	ip authentication key-chain eigrp <i>autonomous-system</i> <i>key-chain</i>	Enable authentication of IP EIGRP packets.
Step 5	exit	Return to global configuration mode.
Step 6	key chain <i>name-of-chain</i>	Identify a key chain and enter key-chain configuration mode. Match the name configured in Step 4.
Step 7	key <i>number</i>	In key-chain configuration mode, identify the key number.
Step 8	key-string <i>text</i>	In key-chain key configuration mode, identify the key string.
Step 9	accept-lifetime <i>start-time</i> { infinite <i>end-time</i> duration <i>seconds</i> }	(Optional) Specify the time period during which the key can be received. The <i>start-time</i> and <i>end-time</i> syntax can be either <i>hh:mm:ss Month date year</i> or <i>hh:mm:ss date Month year</i> . The default is forever with the default <i>start-time</i> and the earliest acceptable date as January 1, 1993. The default <i>end-time</i> and duration is infinite .
Step 10	send-lifetime <i>start-time</i> { infinite <i>end-time</i> duration <i>seconds</i> }	(Optional) Specify the time period during which the key can be sent. The <i>start-time</i> and <i>end-time</i> syntax can be either <i>hh:mm:ss Month date year</i> or <i>hh:mm:ss date Month year</i> . The default is forever with the default <i>start-time</i> and the earliest acceptable date as January 1, 1993. The default <i>end-time</i> and duration is infinite .
Step 11	end	Return to privileged EXEC mode.
Step 12	show key chain	Display authentication key information.
Step 13	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no** forms of these commands to disable the feature or to return the setting to the default value.

Monitoring and Maintaining EIGRP

You can delete neighbors from the neighbor table. You can also display various EIGRP routing statistics. [Table 32-9](#) lists the privileged EXEC commands for deleting neighbors and displaying statistics. For explanations of fields in the resulting display, see the *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*.

Table 32-9 IP EIGRP Clear and Show Commands

Command	Purpose
<code>clear ip eigrp neighbors [if-address interface]</code>	Delete neighbors from the neighbor table.
<code>show ip eigrp interface [interface] [as number]</code>	Display information about interfaces configured for EIGRP.
<code>show ip eigrp neighbors [type-number]</code>	Display EIGRP discovered neighbors.
<code>show ip eigrp topology [autonomous-system-number] [[ip-address] mask]</code>	Display the EIGRP topology table for a given process.
<code>show ip eigrp traffic [autonomous-system-number]</code>	Display the number of packets sent and received for all or a specified EIGRP process.

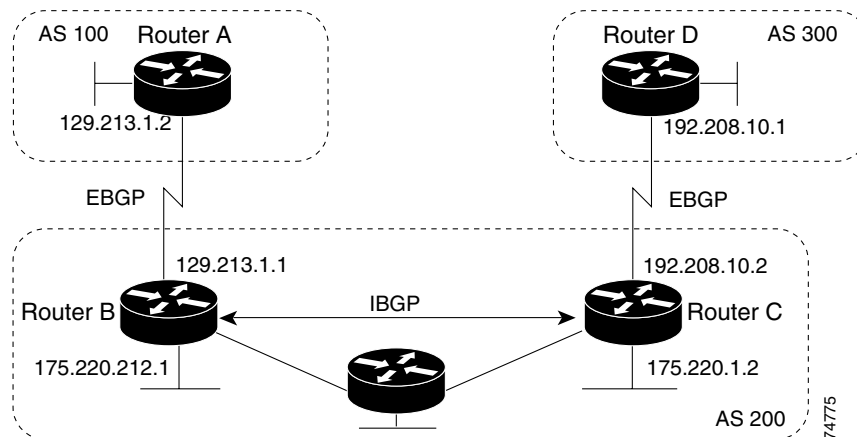
Configuring BGP

The Border Gateway Protocol (BGP) is an exterior gateway protocol used to set up an interdomain routing system that guarantees the loop-free exchange of routing information between autonomous systems. Autonomous systems are made up of routers that operate under the same administration and that run Interior Gateway Protocols (IGPs), such as RIP or OSPF, within their boundaries and that interconnect by using an Exterior Gateway Protocol (EGP). BGP version 4 is the standard EGP for interdomain routing in the Internet. The protocol is defined in RFCs 1163, 1267, and 1771. You can find detailed information about BGP in *Internet Routing Architectures*, published by Cisco Press, and in the “Configuring BGP” chapter in the “IP Routing Protocols” part of the *Cisco IOS IP Configuration Guide, Release 12.2*.

For details about BGP commands and keywords, see the *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*. For a list of BGP commands not supported by the switch, see [Appendix C, “Unsupported CLI Commands in Cisco IOS Release 12.2\(25\)SEE.”](#)

Routers that belong to the same autonomous system (AS) and that exchange BGP updates run *internal BGP* (IBGP), and routers that belong to different autonomous systems and that exchange BGP updates run *external BGP* (EBGP). Most configuration commands are the same for configuring EBGP and IBGP. The difference is that the routing updates are exchanged either between autonomous systems (EBGP) or within an AS (IBGP). [Figure 32-4](#) shows a network that is running both EBGP and IBGP.

Figure 32-4 EBGP, IBGP, and Multiple Autonomous Systems



Before exchanging information with an external AS, BGP ensures that networks within the AS can be reached by defining internal BGP peering among routers within the AS and by redistributing BGP routing information to IGPs that run within the AS, such as IGRP and OSPF.

Routers that run a BGP routing process are often referred to as BGP *speakers*. BGP uses the Transmission Control Protocol (TCP) as its transport protocol (specifically port 179). Two BGP speakers that have a TCP connection to each other for exchanging routing information are known as *peers* or *neighbors*. In [Figure 32-4](#), Routers A and B are BGP peers, as are Routers B and C and Routers C and D. The routing information is a series of AS numbers that describe the full path to the destination network. BGP uses this information to construct a loop-free map of autonomous systems.

The network has these characteristics:

- Routers A and B are running EBGP, and Routers B and C are running IBGP. Note that the EBGP peers are directly connected and that the IBGP peers are not. As long as there is an IGP running that allows the two neighbors to reach one another, IBGP peers do not have to be directly connected.

- All BGP speakers within an AS must establish a peer relationship with each other. That is, the BGP speakers within an AS must be fully meshed logically. BGP4 provides two techniques that reduce the requirement for a logical full mesh: *confederations* and *route reflectors*.
- AS 200 is a *transit AS* for AS 100 and AS 300—that is, AS 200 is used to transfer packets between AS 100 and AS 300.

BGP peers initially exchange their full BGP routing tables and then send only incremental updates. BGP peers also exchange keepalive messages (to ensure that the connection is up) and notification messages (in response to errors or special conditions).

In BGP, each route consists of a network number, a list of autonomous systems that information has passed through (the *autonomous system path*), and a list of other *path attributes*. The primary function of a BGP system is to exchange network reachability information, including information about the list of AS paths, with other BGP systems. This information can be used to determine AS connectivity, to prune routing loops, and to enforce AS-level policy decisions.

A router or switch running Cisco IOS does not select or use an IBGP route unless it has a route available to the next-hop router and it has received synchronization from an IGP (unless IGP synchronization is disabled). When multiple routes are available, BGP bases its path selection on *attribute* values. See the “Configuring BGP Decision Attributes” section on page 32-49 for information about BGP attributes.

BGP Version 4 supports classless interdomain routing (CIDR) so you can reduce the size of your routing tables by creating aggregate routes, resulting in *supernets*. CIDR eliminates the concept of network classes within BGP and supports the advertising of IP prefixes.

These sections briefly describe how to configure BGP and supported BGP features:

- [Default BGP Configuration, page 32-43](#)
- [Enabling BGP Routing, page 32-45](#)
- [Managing Routing Policy Changes, page 32-47](#)
- [Configuring BGP Decision Attributes, page 32-49](#)
- [Configuring BGP Filtering with Route Maps, page 32-51](#)
- [Configuring BGP Filtering by Neighbor, page 32-51](#)
- [Configuring Prefix Lists for BGP Filtering, page 32-53](#)
- [Configuring BGP Community Filtering, page 32-54](#)
- [Configuring BGP Neighbors and Peer Groups, page 32-55](#)
- [Configuring Aggregate Addresses, page 32-57](#)
- [Configuring a Routing Domain Confederation, page 32-58](#)
- [Configuring BGP Route Reflectors, page 32-58](#)
- [Configuring Route Dampening, page 32-59](#)
- [Monitoring and Maintaining BGP, page 32-60](#)

For detailed descriptions of BGP configuration, see the “Configuring BGP” chapter in the “IP Routing Protocols” part of the *Cisco IOS IP Configuration Guide, Release 12.2*. For details about specific commands, see the *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*.

For a list of BGP commands that are visible but not supported by the switch, see [Appendix C, “Unsupported CLI Commands in Cisco IOS Release 12.2\(25\)SEE.”](#)

Default BGP Configuration

Table 32-10 shows the basic default BGP configuration. For the defaults for all characteristics, see the specific commands in the *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*.

Table 32-10 **Default BGP Configuration**

Feature	Default Setting
Aggregate address	Disabled: None defined.
AS path access list	None defined.
Auto summary	Enabled.
Best path	- The router considers <i>as-path</i> in choosing a route and does not compare similar routes from external BGP peers. - Compare router ID: Disabled.
BGP community list	- Number: None defined. When you permit a value for the community number, the list defaults to an implicit deny for everything else that has not been permitted. - Format: Cisco default format (32-bit number).
BGP confederation identifier/peers	- Identifier: None configured. - Peers: None identified.
BGP Fast external fallover	Enabled.
BGP local preference	100. The range is 0 to 4294967295 with the higher value preferred.
BGP network	None specified; no backdoor route advertised.
BGP route dampening	Disabled by default. When enabled: - Half-life is 15 minutes. - Re-use is 750 (10-second increments). - Suppress is 2000 (10-second increments). - Max-suppress-time is 4 times half-life; 60 minutes.
BGP router ID	The IP address of a loopback interface if one is configured or the highest IP address configured for a physical interface on the router.
Default information originate (protocol or network redistribution)	Disabled.
Default metric	Built-in, automatic metric translations.
Distance	- External route administrative distance: 20 (acceptable values are from 1 to 255). - Internal route administrative distance: 200 (acceptable values are from 1 to 255). - Local route administrative distance: 200 (acceptable values are from 1 to 255).
Distribute list	- In (filter networks received in updates): Disabled. - Out (suppress networks from being advertised in updates): Disabled.
Internal route redistribution	Disabled.
IP prefix list	None defined.

Table 32-10 **Default BGP Configuration (continued)**

Feature	Default Setting
Multi exit discriminator (MED)	- Always compare: Disabled. Does not compare MEDs for paths from neighbors in different autonomous systems. - Best path compare: Disabled. - MED missing as worst path: Disabled. - Deterministic MED comparison is disabled.
Neighbor	- Advertisement interval: 30 seconds for external peers; 5 seconds for internal peers. - Change logging: Enabled. - Conditional advertisement: Disabled. - Default originate: No default route is sent to the neighbor. - Description: None. - Distribute list: None defined. - External BGP multihop: Only directly connected neighbors are allowed. - Filter list: None used. - Maximum number of prefixes received: No limit. - Next hop (router as next hop for BGP neighbor): Disabled. - Password: Disabled.
Neighbor	- Peer group: None defined; no members assigned. - Prefix list: None specified. - Remote AS (add entry to neighbor BGP table): No peers defined. - Private AS number removal: Disabled. - Route maps: None applied to a peer. - Send community attributes: None sent to neighbors. - Shutdown or soft reconfiguration: Not enabled. - Timers: keepalive: 60 seconds; holdtime: 180 seconds. - Update source: Best local address. - Version: BGP version 4. - Weight: Routes learned through BGP peer: 0; routes sourced by the local router: 32768.
NSF ¹ Awareness	Disabled ² . Allows Layer 3 switches to continue forwarding packets from a neighboring NSF-capable router during hardware or software changes.
Route reflector	None configured.
Synchronization (BGP and IGP)	Enabled.
Table map update	Disabled.
Timers	Keepalive: 60 seconds; holdtime: 180 seconds.

1. NSF = Nonstop Forwarding

2. NSF Awareness can be enabled on Catalyst 3550 switches with the Cisco IOS Release 12.2(25)SEC IP services image by enabling Graceful Restart.

Nonstop Forwarding Awareness

The BGP NSF awareness feature is supported in the IP services image, beginning with Cisco IOS Release 12.2(25)SEC. To enable this feature with BGP routing, you need to enable Graceful Restart. When the neighboring router is NSF-capable, and this feature is enabled, the Layer 3 switch continues to forward packets from the neighboring router during the interval between the primary Route Processor (RP) in a router failing and the backup RP taking over, or while the primary RP is manually reloaded for a nondisruptive software upgrade.

Disabling Graceful Restart disables NSF awareness.

For more information, see the *BGP Nonstop Forwarding (NSF) Awareness* at this URL:

http://www-search.cisco.com/en/US/products/ps6350/products_configuration_guide_chapter09186a008045568e.html#wp1027129

Enabling BGP Routing

To enable BGP routing, you establish a BGP routing process and define the local network. Because BGP must completely understand the relationships with its neighbors, you must also specify a BGP neighbor.

BGP supports two kinds of neighbors: internal and external. *Internal neighbors* are in the same AS; *external neighbors* are in different autonomous systems. External neighbors are usually adjacent to each other and share a subnet, but internal neighbors can be anywhere in the same AS.

The switch supports the use of private AS numbers, usually assigned by service providers and given to systems whose routes are not advertised to external neighbors. The private AS numbers are from 64512 to 65535. You can configure external neighbors to remove private AS numbers from the AS path by using the **neighbor remove-private-as** router configuration command. Then when an update is passed to an external neighbor, if the AS path includes private AS numbers, these numbers are dropped.

If your AS will be passing traffic through it from another AS to a third AS, it is important to be consistent about the routes it advertises. If BGP advertised a route before all routers in the network had learned about the route through the IGP, the AS might receive traffic that some routers could not yet route. To prevent this from happening, BGP must wait until the IGP has propagated information across the AS so that BGP is *synchronized* with the IGP. Synchronization is enabled by default. If your AS does not pass traffic from one AS to another AS, or if all routers in your autonomous systems are running BGP, you can disable synchronization, which allows your network to carry fewer routes in the IGP and allows BGP to converge more quickly.

Beginning in privileged EXEC mode, follow these steps to enable BGP routing, establish a BGP routing process, and specify a neighbor:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip routing	Enable IP routing (required only if IP routing is disabled).
Step 3	router bgp <i>autonomous-system</i>	Enable a BGP routing process, assign it an AS number, and enter router configuration mode. The AS number can be from 1 to 65535, with 64512 to 65535 designated as private autonomous numbers.
Step 4	network <i>network-number</i> [mask <i>network-mask</i>] [route-map <i>route-map-name</i>]	Configure a network as local to this AS, and enter it in the BGP table.

	Command	Purpose
Step 5	neighbor { <i>ip-address</i> <i>peer-group-name</i> } remote-as <i>number</i>	Add an entry to the BGP neighbor table specifying that the neighbor identified by the IP address belongs to the specified AS. For EBGp, neighbors are usually directly connected, and the IP address is the address of the interface at the other end of the connection. For IBGP, the IP address can be the address of any of the router interfaces.
Step 6	neighbor { <i>ip-address</i> <i>peer-group-name</i> } remove-private-as	(Optional) Remove private AS numbers from the AS-path in outbound routing updates.
Step 7	no synchronization	(Optional) Disable synchronization between BGP and an IGP.
Step 8	no auto-summary	(Optional) Disable automatic network summarization. By default, when a subnet is redistributed from an IGP into BGP, only the network route is inserted into the BGP table.
Step 9	bgp fast-external-fallover	(Optional) Automatically reset a BGP session when a link between external neighbors goes down. By default, the session is not immediately reset.
Step 10	bgp graceful-restart	(Optional) Enable NSF awareness on switch. By default, NSF awareness is disabled.
Step 11	end	Return to privileged EXEC mode.
Step 12	show ip bgp network <i>network-number</i> or show ip bgp neighbor	Verify the configuration. Verify that NSF awareness (Graceful-Restart) is enabled on the neighbor. If NSF awareness is enabled on the switch and the neighbor, this message appears: <i>Graceful Restart Capability: advertised and received</i> If NSF awareness is enabled on the switch, but not on the neighbor, this message appears: <i>Graceful Restart Capability: advertised</i>
Step 13	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no router bgp** *autonomous-system* global configuration command to remove a BGP AS. Use the **no network** *network-number* router configuration command to remove the network from the BGP table. Use the **no neighbor** {*ip-address* | *peer-group-name*} **remote-as** *number* router configuration command to remove a neighbor. Use the **no neighbor** {*ip-address* | *peer-group-name*} **remove-private-as** router configuration command to include private AS numbers in updates to a neighbor. Use the **synchronization** router configuration command to re-enable synchronization.

These examples show how to configure BGP on the routers in [Figure 32-4](#).

Router A:

```
Switch(config)# router bgp 100
Switch(config-router)# neighbor 129.213.1.1 remote-as 200
```

Router B:

```
Switch(config)# router bgp 200
Switch(config-router)# neighbor 129.213.1.2 remote-as 100
Switch(config-router)# neighbor 175.220.1.2 remote-as 200
```

Router C:

```
Switch(config)# router bgp 200
Switch(config-router)# neighbor 175.220.212.1 remote-as 200
Switch(config-router)# neighbor 192.208.10.1 remote-as 300
```

Router D:

```
Switch(config)# router bgp 300
Switch(config-router)# neighbor 192.208.10.2 remote-as 200
```

To verify that BGP peers are running, use the **show ip bgp neighbors** privileged EXEC command. This is the output of this command on Router A:

```
Switch# show ip bgp neighbors

BGP neighbor is 129.213.1.1, remote AS 200, external link
  BGP version 4, remote router ID 175.220.212.1
  BGP state = established, table version = 3, up for 0:10:59
  Last read 0:00:29, hold time is 180, keepalive interval is 60 seconds
  Minimum time between advertisement runs is 30 seconds
  Received 2828 messages, 0 notifications, 0 in queue
  Sent 2826 messages, 0 notifications, 0 in queue
  Connections established 11; dropped 10
```

Anything other than *state = established* means that the peers are not running. The remote router ID is the highest IP address on that router (or the highest loopback interface). Each time the table is updated with new information, the table version number increments. A table version number that continually increments means that a route is flapping, causing continual routing updates.

For exterior protocols, a reference to an IP network from the **network** router configuration command controls only which networks are advertised. This is in contrast to Interior Gateway Protocols (IGPs), such as IGRP, which also use the **network** command to determine where to send updates.

For detailed descriptions of BGP configuration, see the “IP Routing Protocols” part in the *Cisco IOS IP Configuration Guide, Release 12.2*. For details about specific commands, see the *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*. For a list of BGP commands that are visible but not supported by the switch, see [Appendix C, “Unsupported CLI Commands in Cisco IOS Release 12.2\(25\)SEE.”](#)

Managing Routing Policy Changes

Routing policies for a peer include all the configurations that might affect inbound or outbound routing table updates. When you have defined two routers as BGP neighbors, they form a BGP connection and exchange routing information. If you later change a BGP filter, weight, distance, version, or timer, or make a similar configuration change, you must reset the BGP sessions so that the configuration changes take effect.

There are two types of reset, hard reset and soft reset. Cisco IOS software releases 12.1 and later support a soft reset without any prior configuration. To use a soft reset without preconfiguration, both BGP peers must support the soft route refresh capability, which is advertised in the OPEN message sent when the peers establish a TCP session. A soft reset allows the dynamic exchange of route refresh requests and routing information between BGP routers and the subsequent re-advertisement of the respective outbound routing table.

- When soft reset generates inbound updates from a neighbor, it is called dynamic inbound soft reset.
- When soft reset sends a set of updates to a neighbor, it is called outbound soft reset.

A soft inbound reset enables the new inbound policy to take effect. A soft outbound reset causes the new local outbound policy to take effect without resetting the BGP session. As a new set of updates is sent during outbound policy reset, a new inbound policy can also take effect.

Table 32-11 lists the advantages and disadvantages hard reset and soft reset.

Table 32-11 Advantages and Disadvantages of Hard and Soft Resets

Type of Reset	Advantages	Disadvantages
hard reset	No memory overhead.	The prefixes in the BGP, IP, and FIB tables provided by the neighbor are lost. Not recommended.
outbound soft reset	No configuration, no storing of routing table updates.	Does not reset inbound routing table updates.
dynamic inbound soft reset	Does not clear the BGP session and cache. Does not require storing of routing table updates and has no memory overhead.	Both BGP routers must support the route refresh capability (in Cisco IOS Release 12.1 and later releases).

Beginning in privileged EXEC mode, follow these steps to determine if a BGP peer supports the route refresh capability and to reset the BGP session:

	Command	Purpose
Step 1	show ip bgp neighbors	Display whether a neighbor supports the route refresh capability. When supported, this message appears for the router: <i>Received route refresh capability from peer.</i>
Step 2	clear ip bgp { * <i>address</i> <i>peer-group-name</i> }	Reset the routing table on the specified connection. • Enter an asterisk (*) to specify that all connections be reset. • Enter an IP <i>address</i> to specify the connection to be reset. • Enter a peer group name to reset the peer group.
Step 3	clear ip bgp { * <i>address</i> <i>peer-group-name</i> } soft out	(Optional) Perform an outbound soft reset to reset the inbound routing table on the specified connection. Use this command if route refresh is supported. • Enter an asterisk (*) to specify that all connections be reset. • Enter an IP <i>address</i> to specify the connection to be reset. • Enter a peer group name to reset the peer group.
Step 4	show ip bgp show ip bgp neighbors	Verify the reset by checking information about the routing table and about BGP neighbors.

Configuring BGP Decision Attributes

When a BGP speaker receives updates from multiple autonomous systems that describe different paths to the same destination, it must choose the single best path for reaching that destination. When chosen, the selected path is entered into the BGP routing table and propagated to its neighbors. The decision is based on the value of attributes that the update contains and other BGP-configurable factors.

When a BGP peer learns two EBGp paths for a prefix from a neighboring AS, it chooses the best path and inserts that path in the IP routing table. If BGP multipath support is enabled and the EBGp paths are learned from the same neighboring autonomous systems, instead of a single best path, multiple paths are installed in the IP routing table. Then, during packet switching, per-packet or per-destination load balancing is performed among the multiple paths. The **maximum-paths** router configuration command controls the number of paths allowed.

These factors summarize the order in which BGP evaluates the attributes for choosing the best path:

1. If the path specifies a next hop that is inaccessible, drop the update. The BGP next-hop attribute, automatically determined by the software, is the IP address of the next hop that is going to be used to reach a destination. For EBGp, this is usually the IP address of the neighbor specified by the **neighbor remote-as** router configuration command. You can disable next-hop processing by using route maps or the **neighbor next-hop-self** router configuration command.
2. Prefer the path with the largest weight (a Cisco proprietary parameter). The weight attribute is local to the router and not propagated in routing updates. By default, the weight attribute is 32768 for paths that the router originates and zero for other paths. Routes with the largest weight are preferred. You can use access lists, route maps, or the **neighbor weight** router configuration command to set weights.
3. Prefer the route with the highest local preference. Local preference is part of the routing update and exchanged among routers in the same AS. The default value of the local preference attribute is 100. You can set local preference by using the **bgp default local-preference** router configuration command or by using a route map.
4. Prefer the route that was originated by BGP running on the local router.
5. Prefer the route with the shortest AS path.
6. Prefer the route with the lowest origin type. An interior route or IGP is lower than a route learned by EGP, and an EGP-learned route is lower than one of unknown origin or learned in another way.
7. Prefer the route with the lowest Multi Exit Discriminator (MED) metric attribute if the neighboring AS is the same for all routes considered. You can configure the MED by using route maps or by using the **default-metric** router configuration command. When an update is sent to an IBGP peer, the MED is included.
8. Prefer the external (EBGP) path over the internal (IBGP) path.
9. Prefer the route that can be reached through the closest IGP neighbor (the lowest IGP metric). This means that the router will prefer the shortest internal path within the AS to reach the destination (the shortest path to the BGP next-hop).
10. If the following conditions are all true, insert the route for this path into the IP routing table:
 - Both the best route and this route are external.
 - Both the best route and this route are from the same neighboring autonomous system.
 - **maximum-paths** is enabled.
11. If multipath is not enabled, prefer the route with the lowest IP address value for the BGP router ID. The router ID is usually the highest IP address on the router or the loopback (virtual) address, but might be implementation-specific.

Beginning in privileged EXEC mode, follow these steps to configure some decision attributes:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router bgp <i>autonomous-system</i>	Enable a BGP routing process, assign it an AS number, and enter router configuration mode.
Step 3	bgp best-path as-path ignore	(Optional) Configure the router to ignore AS path length in selecting a route.
Step 4	neighbor { <i>ip-address</i> <i>peer-group-name</i> } next-hop-self	(Optional) Disable next-hop processing on BGP updates to a neighbor by entering a specific IP address to be used instead of the next-hop address.
Step 5	neighbor { <i>ip-address</i> <i>peer-group-name</i> } weight <i>weight</i>	(Optional) Assign a weight to a neighbor connection. Acceptable values are from 0 to 65535; the largest weight is the preferred route. Routes learned through another BGP peer have a default weight of 0; and routes sourced by the local router have a default weight of 32768.
Step 6	default-metric <i>number</i>	(Optional) Set a MED metric to set preferred paths to external neighbors. All routes without a MED will also be set to this value. The range is 1 to 4294967295. The lowest value is the most desirable.
Step 7	bgp bestpath med missing-as-worst	(Optional) Configure the switch to consider a missing MED as having a value of infinity, making the path without a MED value the least desirable path.
Step 8	bgp always-compare med	(Optional) Configure the switch to compare MEDs for paths from neighbors in different autonomous systems. By default, MED comparison is only done among paths in the same AS.
Step 9	bgp bestpath med confed	(Optional) Configure the switch to consider the MED in choosing a path from among those advertised by different subautonomous systems within a confederation.
Step 10	bgp deterministic med	(Optional) Configure the switch to consider the MED variable when choosing among routes advertised by different peers in the same AS.
Step 11	bgp default local-preference <i>value</i>	(Optional) Change the default local preference value. The range is 0 to 4294967295; the default value is 100. The highest local preference value is preferred.
Step 12	maximum-paths <i>number</i>	(Optional) Configure the number of paths to be added to the IP routing table. The default is to only enter the best path in the routing table. The range is from 1 to 8. Having multiple paths allows load balancing among the paths. (Although the switch software allows a maximum of 32 equal-cost routes, the switch hardware will never use more than 8 paths per route.)
Step 13	end	Return to privileged EXEC mode.

	Command	Purpose
Step 14	show ip bgp show ip bgp neighbors	Verify the reset by checking information about the routing table and about BGP neighbors.
Step 15	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no** form of each command to return to the default state.

Configuring BGP Filtering with Route Maps

Within BGP, route maps can be used to control and to modify routing information and to define the conditions by which routes are redistributed between routing domains. See the [“Using Route Maps to Redistribute Routing Information” section on page 32-76](#) for more information about route maps. Each route map has a name that identifies the route map (*map tag*) and an optional sequence number.

Beginning in privileged EXEC mode, follow these steps to use a route map to disable next-hop processing:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	route-map <i>map-tag</i> [[permit deny] <i>sequence-number</i>]]	Create a route map, and enter route-map configuration mode.
Step 3	set ip next-hop <i>ip-address</i> [... <i>ip-address</i>] [<i>peer-address</i>]	(Optional) Set a route map to disable next-hop processing - In an inbound route map, set the next hop of matching routes to be the neighbor peering address, overriding third-party next hops. - In an outbound route map of a BGP peer, set the next hop to the peering address of the local router, disabling the next-hop calculation.
Step 4	end	Return to privileged EXEC mode.
Step 5	show route-map [<i>map-name</i>]	Display all route maps configured or only the one specified to verify configuration.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no route-map** *map-tag* command to delete the route map. Use the **no set ip next-hop** *ip-address* command to re-enable next-hop processing.

Configuring BGP Filtering by Neighbor

You can filter BGP advertisements by using AS-path filters, such as the **as-path access-list** global configuration command and the **neighbor filter-list** router configuration command. You can also use access lists with the **neighbor distribute-list** router configuration command. Distribute-list filters are applied to network numbers. See the [“Controlling Advertising and Processing in Routing Updates” section on page 32-83](#) for information about the **distribute-list** command.

You can use route maps on a per-neighbor basis to filter updates and to modify various attributes. A route map can be applied to either inbound or outbound updates. Only the routes that pass the route map are sent or accepted in updates. On both inbound and outbound updates, matching is supported based on AS

path, community, and network numbers. Autonomous system path matching requires the **match as-path access-list** route-map command, community based matching requires the **match community-list** route-map command, and network-based matching requires the **ip access-list** global configuration command.

Beginning in privileged EXEC mode, follow these steps to apply a per-neighbor route map:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router bgp <i>autonomous-system</i>	Enable a BGP routing process, assign it an AS number, and enter router configuration mode.
Step 3	neighbor { <i>ip-address</i> <i>peer-group name</i> } distribute-list { <i>access-list-number</i> <i>name</i> } { in out }	(Optional) Filter BGP routing updates to or from neighbors as specified in an access list. Note You can also use the neighbor prefix-list router configuration command to filter updates, but you cannot use both commands to configure the same BGP peer.
Step 4	neighbor { <i>ip-address</i> <i>peer-group name</i> } route-map <i>map-tag</i> { in out }	(Optional) Apply a route map to filter an incoming or outgoing route.
Step 5	end	Return to privileged EXEC mode.
Step 6	show ip bgp neighbors	Verify the configuration.
Step 7	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no neighbor distribute-list** command to remove the access list from the neighbor. Use the **no neighbor route-map map-tag** router configuration command to remove the route map from the neighbor.

Another method of filtering is to specify an access list filter on both incoming and outbound updates, based on the BGP autonomous system paths. Each filter is an access list based on regular expressions. (See the “Regular Expressions” appendix in the *Cisco IOS Dial Technologies Command Reference, Release 12.1* for more information on forming regular expressions.) To use this method, define an autonomous system path access list, and apply it to updates to and from particular neighbors.

Beginning in privileged EXEC mode, follow these steps to configure BGP path filtering:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip as-path access-list <i>access-list-number</i> { permit deny } <i>as-regular-expressions</i>	Define a BGP-related access list.
Step 3	router bgp <i>autonomous-system</i>	Enter BGP router configuration mode.
Step 4	neighbor { <i>ip-address</i> <i>peer-group name</i> } filter-list { <i>access-list-number</i> <i>name</i> } { in out weight <i>weight</i> }	Establish a BGP filter based on an access list.
Step 5	end	Return to privileged EXEC mode.
Step 6	show ip bgp neighbors [paths regular-expression]	Verify the configuration.
Step 7	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Configuring Prefix Lists for BGP Filtering

You can use prefix lists as an alternative to access lists in many BGP route filtering commands, including the **neighbor distribute-list** router configuration command. The advantages of using prefix lists include performance improvements in loading and lookup of large lists, incremental update support, easier CLI configuration, and greater flexibility.

Filtering by a prefix list involves matching the prefixes of routes with those listed in the prefix list, as when matching access lists. When there is a match, the route is used. Whether a prefix is permitted or denied is based upon these rules:

- An empty prefix list permits all prefixes.
- An implicit deny is assumed if a given prefix does not match any entries in a prefix list.
- When multiple entries of a prefix list match a given prefix, the sequence number of a prefix list entry identifies the entry with the lowest sequence number.

By default, sequence numbers are generated automatically and incremented in units of five. If you disable the automatic generation of sequence numbers, you must specify the sequence number for each entry. You can specify sequence values in any increment. If you specify increments of one, you cannot insert additional entries into the list; if you choose very large increments, you might run out of values.

You do not need to specify a sequence number when removing a configuration entry. **Show** commands include the sequence numbers in their output.

Before using a prefix list in a command, you must set up the prefix list. Beginning in privileged EXEC mode, follow these steps to create a prefix list or to add an entry to a prefix list:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip prefix-list <i>list-name</i> [seq <i>seq-value</i>] deny permit <i>network/len</i> [ge <i>ge-value</i>] [le <i>le-value</i>]	Create a prefix list with an optional sequence number to deny or permit access for matching conditions. You must enter at least one permit or deny clause. • <i>network/len</i> is the network number and length (in bits) of the network mask. • (Optional) ge and le values specify the range of the prefix length to be matched. The specified <i>ge-value</i> and <i>le-value</i> must satisfy this condition: $len < ge\text{-}value < le\text{-}value < 32$
Step 3	ip prefix-list <i>list-name</i> seq <i>seq-value</i> deny permit <i>network/len</i> [ge <i>ge-value</i>] [le <i>le-value</i>]	(Optional) Add an entry to a prefix list, and assign a sequence number to the entry.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip prefix list [detail summary] <i>name</i> [<i>network/len</i>] [seq <i>seq-num</i>] [longer] [first-match]	Verify the configuration by displaying information about a prefix list or prefix list entries.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To delete a prefix list and all of its entries, use the **no ip prefix-list** *list-name* global configuration command. To delete an entry from a prefix list, use the **no ip prefix-list seq** *seq-value* global configuration command. To disable automatic generation of sequence numbers, use the **no ip prefix-list**

sequence number command; to reenale automatic generation, use the **ip prefix-list sequence number** command. To clear the hit-count table of prefix list entries, use the **clear ip prefix-list** privileged EXEC command.

Configuring BGP Community Filtering

One way that BGP controls the distribution of routing information based on the value of the COMMUNITIES attribute. The attribute is a way to groups destinations into communities and to apply routing decisions based on the communities. This method simplifies configuration of a BGP speaker to control distribution of routing information.

A *community* is a group of destinations that share some common attribute. Each destination can belong to multiple communities. AS administrators can define to which communities a destination belongs. By default, all destinations belong to the general Internet community. The community is identified by the COMMUNITIES attribute, an optional, transitive, global attribute in the numerical range from 1 to 4294967200. These are some predefined, well-known communities:

- **internet**—Advertise this route to the Internet community. All routers belong to it.
- **no-export**—Do not advertise this route to EBGp peers.
- **no-advertise**—Do not advertise this route to any peer (internal or external).
- **local-as**—Do not advertise this route to peers outside the local autonomous system.

Based on the community, you can control which routing information to accept, prefer, or distribute to other neighbors. A BGP speaker can set, append, or modify the community of a route when learning, advertising, or redistributing routes. When routes are aggregated, the resulting aggregate has a COMMUNITIES attribute that contains all communities from all the initial routes.

You can use community lists to create groups of communities to use in a match clause of a route map. As with an access list, a series of community lists can be created. Statements are checked until a match is found. As soon as one statement is satisfied, the test is concluded.

To set the COMMUNITIES attribute and match clauses based on communities, see the **match community-list** and **set community** route-map configuration commands in the “Using Route Maps to Redistribute Routing Information” section on page 32-76.

By default, no COMMUNITIES attribute is sent to a neighbor. You can specify that the COMMUNITIES attribute be sent to the neighbor at an IP address by using the **neighbor send-community** router configuration command.

Beginning in privileged EXEC mode, follow these steps to create and to apply a community list:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip community-list <i>community-list-number</i> { permit deny } <i>community-number</i>	Create a community list and assign it a number. • The <i>community-list-number</i> is an integer from 1 to 99 that identifies one or more permit or deny groups of communities. • The <i>community-number</i> is the number configured by a set community route-map configuration command.
Step 3	router bgp <i>autonomous-system</i>	Enter BGP router configuration mode.
Step 4	neighbor { <i>ip-address</i> <i>peer-group name</i> } send-community	Specify that the COMMUNITIES attribute be sent to the neighbor at this IP address.

	Command	Purpose
Step 5	set comm-list <i>list-num</i> delete	(Optional) Remove communities from the community attribute of an inbound or outbound update that match a standard or extended community list specified by a route map.
Step 6	exit	Return to global configuration mode.
Step 7	ip bgp-community new-format	(Optional) Display and parse BGP communities in the format AA:NN. A BGP community appears in a 2-part format two bytes long. The Cisco default community format is in the format NNAA. In the most recent RFC for BGP, a community takes the form AA:NN, where the first part is the AS number and the second part is a 2-byte number.
Step 8	end	Return to privileged EXEC mode.
Step 9	show ip bgp community	Verify the configuration.
Step 10	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Configuring BGP Neighbors and Peer Groups

Often many BGP neighbors are configured with the same update policies (that is, the same outbound route maps, distribute lists, filter lists, update source, and so on). Neighbors with the same update policies can be grouped into peer groups to simplify configuration and to make updating more efficient. When you have configured many peers, we recommend this approach.

To configure a BGP peer group, you create the peer group, assign options to the peer group, and add neighbors as peer group members. You configure the peer group by using the **neighbor** router configuration commands. By default, peer group members inherit all the configuration options of the peer group, including the remote-as (if configured), version, update-source, out-route-map, out-filter-list, out-dist-list, minimum-advertisement-interval, and next-hop-self. All peer group members also inherit changes made to the peer group. Members can also be configured to override the options that do not affect outbound updates.

To assign configuration options to an individual neighbor, specify any of these router configuration commands by using the neighbor IP address. To assign the options to a peer group, specify any of the commands by using the peer group name. You can disable a BGP peer or peer group without removing all the configuration information by using the **neighbor shutdown** router configuration command.

Beginning in privileged EXEC mode, use these commands to configure BGP peers:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router bgp <i>autonomous-system</i>	Enter BGP router configuration mode.
Step 3	neighbor <i>peer-group-name</i> peer-group	Create a BGP peer group.
Step 4	neighbor <i>ip-address</i> peer-group <i>peer-group-name</i>	Make a BGP neighbor a member of the peer group.
Step 5	neighbor { <i>ip-address</i> <i>peer-group-name</i> } remote-as <i>number</i>	Specify a BGP neighbor. If a peer group is not configured with a remote-as <i>number</i> , use this command to create peer groups containing EBGP neighbors. The range is 1 to 65535.
Step 6	neighbor { <i>ip-address</i> <i>peer-group-name</i> } description <i>text</i>	(Optional) Associate a description with a neighbor.

	Command	Purpose
Step 7	neighbor { <i>ip-address</i> <i>peer-group-name</i> } default-originate [route-map <i>map-name</i>]	(Optional) Allow a BGP speaker (the local router) to send the default route 0.0.0.0 to a neighbor for use as a default route.
Step 8	neighbor { <i>ip-address</i> <i>peer-group-name</i> } send-community	(Optional) Specify that the COMMUNITIES attribute be sent to the neighbor at this IP address.
Step 9	neighbor { <i>ip-address</i> <i>peer-group-name</i> } update-source <i>interface</i>	(Optional) Allow internal BGP sessions to use any operational interface for TCP connections.
Step 10	neighbor { <i>ip-address</i> <i>peer-group-name</i> } ebgp-multihop	(Optional) Allow BGP sessions, even when the neighbor is not on a directly connected segment. The multihop session is not established if the only route to the multihop peer's address is the default route (0.0.0.0).
Step 11	neighbor { <i>ip-address</i> <i>peer-group-name</i> } local-as <i>number</i>	(Optional) Specify an AS number to use as the local AS. The range is 1 to 65535.
Step 12	neighbor { <i>ip-address</i> <i>peer-group-name</i> } advertisement-interval <i>seconds</i>	(Optional) Set the minimum interval between sending BGP routing updates.
Step 13	neighbor { <i>ip-address</i> <i>peer-group-name</i> } maximum-prefix <i>maximum</i> [<i>threshold</i>]	(Optional) Control how many prefixes can be received from a neighbor. The range is 1 to 4294967295. The <i>threshold</i> (optional) is the percentage of maximum at which a warning message is generated. The default is 75 percent.
Step 14	neighbor { <i>ip-address</i> <i>peer-group-name</i> } next-hop-self	(Optional) Disable next-hop processing on the BGP updates to a neighbor.
Step 15	neighbor { <i>ip-address</i> <i>peer-group-name</i> } password <i>string</i>	(Optional) Set MD5 authentication on a TCP connection to a BGP peer. The same password must be configured on both BGP peers, or the connection between them is not made.
Step 16	neighbor { <i>ip-address</i> <i>peer-group-name</i> } route-map <i>map-name</i> { in out }	(Optional) Apply a route map to incoming or outgoing routes.
Step 17	neighbor { <i>ip-address</i> <i>peer-group-name</i> } send-community	(Optional) Specify that the COMMUNITIES attribute be sent to the neighbor at this IP address.
Step 18	neighbor { <i>ip-address</i> <i>peer-group-name</i> } timers <i>keepalive holdtime</i>	(Optional) Set timers for the neighbor or peer group. - The <i>keepalive</i> interval is the time within which keepalive messages are sent to peers. The range is 1 to 4294967295 seconds; the default is 60. - The <i>holdtime</i> is the interval after which a peer is declared inactive after not receiving a keepalive message from it. The range is 1 to 4294967295 seconds; the default is 180.
Step 19	neighbor { <i>ip-address</i> <i>peer-group-name</i> } weight <i>weight</i>	(Optional) Specify a weight for all routes from a neighbor.
Step 20	neighbor { <i>ip-address</i> <i>peer-group-name</i> } distribute-list { <i>access-list-number</i> <i>name</i> } { in out }	(Optional) Filter BGP routing updates to or from neighbors, as specified in an access list.
Step 21	neighbor { <i>ip-address</i> <i>peer-group-name</i> } filter-list <i>access-list-number</i> { in out weight <i>weight</i> }	(Optional) Establish a BGP filter.
Step 22	neighbor { <i>ip-address</i> <i>peer-group-name</i> } version <i>value</i>	(Optional) Specify the BGP version to use when communicating with a neighbor.

	Command	Purpose
Step 23	neighbor { <i>ip-address</i> <i>peer-group-name</i> } soft-reconfiguration inbound	(Optional) Configure the software to start storing received updates.
Step 24	end	Return to privileged EXEC mode.
Step 25	show ip bgp neighbors	Verify the configuration.
Step 26	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To disable an existing BGP neighbor or neighbor peer group, use the **neighbor shutdown** router configuration command. To enable a previously existing neighbor or neighbor peer group that had been disabled, use the **no neighbor shutdown** router configuration command.

Configuring Aggregate Addresses

Classless interdomain routing (CIDR) enables you to create aggregate routes (or *supernets*) to minimize the size of routing tables. You can configure aggregate routes in BGP either by redistributing an aggregate route into BGP or by creating an aggregate entry in the BGP routing table. An aggregate address is added to the BGP table when there is at least one more specific entry in the BGP table.

Beginning in privileged EXEC mode, use these commands to create an aggregate address in the routing table:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router bgp <i>autonomous-system</i>	Enter BGP router configuration mode.
Step 3	aggregate-address <i>address mask</i>	Create an aggregate entry in the BGP routing table. The aggregate route is advertised as coming from the AS, and the atomic aggregate attribute is set to indicate that information might be missing.
Step 4	aggregate-address <i>address mask as-set</i>	(Optional) Generate AS set path information. This command creates an aggregate entry following the same rules as the previous command, but the advertised path will be an AS_SET consisting of all elements contained in all paths. Do not use this keyword when aggregating many paths because this route must be continually withdrawn and updated.
Step 5	aggregate-address <i>address-mask summary-only</i>	(Optional) Advertise summary addresses only.
Step 6	aggregate-address <i>address mask suppress-map</i> <i>map-name</i>	(Optional) Suppress selected, more specific routes.
Step 7	aggregate-address <i>address mask advertise-map</i> <i>map-name</i>	(Optional) Generate an aggregate based on conditions specified by the route map.
Step 8	aggregate-address <i>address mask attribute-map</i> <i>map-name</i>	(Optional) Generate an aggregate with attributes specified in the route map.
Step 9	end	Return to privileged EXEC mode.
Step 10	show ip bgp neighbors [advertised-routes]	Verify the configuration.
Step 11	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To delete an aggregate entry, use the **no aggregate-address** *address mask* router configuration command. To return options to the default values, use the command with keywords.

Configuring a Routing Domain Confederation

One way to reduce the IBGP mesh is to divide an autonomous system into multiple subautonomous systems and to group them into a single confederation that appears as a single autonomous system. Each autonomous system is fully meshed within itself and has a few connections to other autonomous systems in the same confederation. Even though the peers in different autonomous systems have EBGP sessions, they exchange routing information as if they were IBGP peers. Specifically, the next hop, MED, and local preference information is preserved. You can then use a single IGP for all of the autonomous systems.

To configure a BGP confederation, you must specify a confederation identifier that acts as the autonomous system number for the group of autonomous systems.

Beginning in privileged EXEC mode, use these commands to configure a BGP confederation:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router bgp <i>autonomous-system</i>	Enter BGP router configuration mode.
Step 3	bgp confederation identifier <i>autonomous-system</i>	Configure a BGP confederation identifier.
Step 4	bgp confederation peers <i>autonomous-system</i> [<i>autonomous-system ...</i>]	Specify the autonomous systems that belong to the confederation and that will be treated as special EBGP peers.
Step 5	end	Return to privileged EXEC mode.
Step 6	show ip bgp neighbor show ip bgp network	Verify the configuration.
Step 7	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Configuring BGP Route Reflectors

BGP requires that all of the IBGP speakers be fully meshed. When a router receives a route from an external neighbor, it must advertise it to all internal neighbors. To prevent a routing information loop, all IBGP speakers must be connected. The internal neighbors do not send routes learned from internal neighbors to other internal neighbors.

With route reflectors, all IBGP speakers need not be fully meshed because another method is used to pass learned routes to neighbors. When you configure an internal BGP peer to be a *route reflector*, it is responsible for passing IBGP learned routes to a set of IBGP neighbors. The internal peers of the route reflector are divided into two groups: *client peers* and *nonclient peers* (all the other routers in the autonomous system). A route reflector reflects routes between these two groups. The route reflector and its client peers form a *cluster*. The nonclient peers must be fully meshed with each other, but the client peers need not be fully meshed. The clients in the cluster do not communicate with IBGP speakers outside their cluster.

When the route reflector receives an advertised route, it takes one of these actions, depending on the neighbor:

- A route from an external BGP speaker is advertised to all clients and nonclient peers.

- A route from a nonclient peer is advertised to all clients.
- A route from a client is advertised to all clients and nonclient peers. Hence, the clients need not be fully meshed.

Usually a cluster of clients have a single route reflector, and the cluster is identified by the route reflector router ID. To increase redundancy and to avoid a single point of failure, a cluster might have more than one route reflector. In this case, all route reflectors in the cluster must be configured with the same 4-byte cluster ID so that a route reflector can recognize updates from route reflectors in the same cluster. All the route reflectors serving a cluster should be fully meshed and should have identical sets of client and nonclient peers.

Beginning in privileged EXEC mode, use these commands to configure a route reflector and clients:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router bgp <i>autonomous-system</i>	Enter BGP router configuration mode.
Step 3	neighbor <i>ip-address</i> <i>peer-group-name</i> route-reflector-client	Configure the local router as a BGP route reflector and the specified neighbor as a client.
Step 4	bgp cluster-id <i>cluster-id</i>	(Optional) Configure the cluster ID if the cluster has more than one route reflector.
Step 5	no bgp client-to-client reflection	(Optional) Disable client-to-client route reflection. By default, the routes from a route reflector client are reflected to other clients. However, if the clients are fully meshed, the route reflector does not need to reflect routes to clients.
Step 6	end	Return to privileged EXEC mode.
Step 7	show ip bgp	Verify the configuration. Display the originator ID and the cluster-list attributes.
Step 8	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Configuring Route Dampening

Route flap dampening is a BGP feature designed to minimize the propagation of flapping routes across an internetwork. A route is considered to be flapping when it is repeatedly available, then unavailable, then available, then unavailable, and so on. When route dampening is enabled, a numeric *penalty* value is assigned to a route when it flaps. When a route's accumulated penalties reach a configurable limit, BGP suppresses advertisements of the route, even if the route is running. The *reuse limit* is a configurable value that is compared with the penalty. If the penalty is less than the reuse limit, a suppressed route that is up will be advertised again.

Dampening is not applied to routes that are learned by IBGP. This policy prevents the IBGP peers from having a higher penalty for routes external to the AS.

Beginning in privileged EXEC mode, use these commands to configure BGP route dampening:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router bgp <i>autonomous-system</i>	Enter BGP router configuration mode.
Step 3	bgp dampening	Enable BGP route dampening.

	Command	Purpose
Step 4	bgp dampening <i>half-life reuse suppress max-suppress [route-map map]</i>	(Optional) Change the default values of route dampening factors.
Step 5	end	Return to privileged EXEC mode.
Step 6	show ip bgp flap-statistics [{ regexp <i>regexp</i> } { filter-list <i>list</i> } { <i>address mask</i> [longer-prefix]}]	(Optional) Monitor the flaps of all paths that are flapping. The statistics are deleted once the route is not suppressed and is stable.
Step 7	show ip bgp dampened-paths	(Optional) Display the dampened routes, including the time remaining before they are suppressed.
Step 8	clear ip bgp flap-statistics [{ regexp <i>regexp</i> } { filter-list <i>list</i> } { <i>address mask</i> [longer-prefix]}]	(Optional) Clear BGP flap statistics to make it less likely that a route will be dampened.
Step 9	clear ip bgp dampening	(Optional) Clear route dampening information and unsuppress the suppressed routes.
Step 10	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To disable flap dampening, use the **no bgp dampening** router configuration command without keywords. To set dampening factors back to the default values, use the **no bgp dampening** router configuration command with values.

Monitoring and Maintaining BGP

You can remove all contents of a particular cache, table, or database. This might be necessary when the contents of the particular structure have become or are suspected to be invalid.

You can display specific statistics, such as the contents of BGP routing tables, caches, and databases. You can use the information to determine resource utilization and solve network problems. You can also display information about node reachability and discover the routing path your device's packets are taking through the network.

Table 32-9 lists the privileged EXEC commands for clearing and displaying BGP. For explanations of the display fields, see the *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*.

Table 32-12 IP BGP Clear and Show Commands

Command	Purpose
clear ip bgp <i>address</i>	Reset a particular BGP connection.
clear ip bgp *	Reset all BGP connections.
clear ip bgp peer-group <i>tag</i>	Remove all members of a BGP peer group.
show ip bgp <i>prefix</i>	Display peer groups and peers not in peer groups to which the prefix has been advertised. Also displays prefix attributes such as the next hop and the local prefix.
show ip bgp cidr-only	Display all BGP routes that contain subnet and supernet network masks.
show ip bgp community [<i>community-number</i>] [exact]	Display routes that belong to the specified communities.
show ip bgp community-list <i>community-list-number</i> [exact-match]	Display routes that are permitted by the community list.

Table 32-12 IP BGP Clear and Show Commands (continued)

Command	Purpose
show ip bgp filter-list <i>access-list-number</i>	Display routes that are matched by the specified AS path access list.
show ip bgp inconsistent-as	Display the routes with inconsistent originating autonomous systems.
show ip bgp regexp <i>regular-expression</i>	Display the routes that have an AS path that matches the specified regular expression entered on the command line.
show ip bgp	Display the contents of the BGP routing table.
show ip bgp neighbors [<i>address</i>]	Display detailed information on the BGP and TCP connections to individual neighbors.
show ip bgp neighbors [<i>address</i>] [advertised-routes dampened-routes flap-statistics paths <i>regular-expression</i> received-routes routes]	Display routes learned from a particular BGP neighbor.
show ip bgp paths	Display all BGP paths in the database.
show ip bgp peer-group [<i>tag</i>] [summary]	Display information about BGP peer groups.
show ip bgp summary	Display the status of all BGP connections.

You can also enable the logging of messages generated when a BGP neighbor resets, comes up, or goes down, by using the **bgp log-neighbor changes** router configuration command.

Configuring Multi-VRF CE

Virtual Private Networks (VPNs) provide a secure way for customers to share bandwidth over an ISP backbone network. A VPN is a collection of sites sharing a common routing table. A customer site is connected to the service-provider network by one or more interfaces, and the service provider associates each interface with a VPN routing table, called a VPN routing/forwarding (VRF) table.

The switch supports multiple VPN routing/forwarding (multi-VRF) instances in customer edge (CE) devices (multi-VRF CE). Multi-VRF CE allows a service provider to support two or more VPNs with overlapping IP addresses.



Note

The switch does not use Multiprotocol Label Switching (MPLS) to support VPNs. For information about MPLS VRF, see the *Cisco IOS Switching Services Configuration Guide, Release 12.2*.

This section includes these topics:

- [Understanding Multi-VRF CE, page 32-62](#)
- [Default Multi-VRF CE Configuration, page 32-64](#)
- [Multi-VRF CE Configuration Guidelines, page 32-65](#)
- [Configuring VRFs, page 32-66](#)
- [Configuring a VPN Routing Session, page 32-67](#)
- [Configuring BGP PE to CE Routing Sessions, page 32-67](#)
- [Multi-VRF CE Configuration Example, page 32-68](#)

- [Displaying Multi-VRF CE Status, page 32-72](#)

Understanding Multi-VRF CE

Multi-VRF CE is a feature that allows a service provider to support two or more VPNs, where IP addresses can be overlapped among the VPNs. Multi-VRF CE uses input interfaces to distinguish routes for different VPNs and forms virtual packet-forwarding tables by associating one or more Layer 3 interfaces with each VRF. Interfaces in a VRF can be either physical, such as Ethernet ports, or logical, such as VLAN SVIs, but an interface cannot belong to more than one VRF at any time.

**Note**

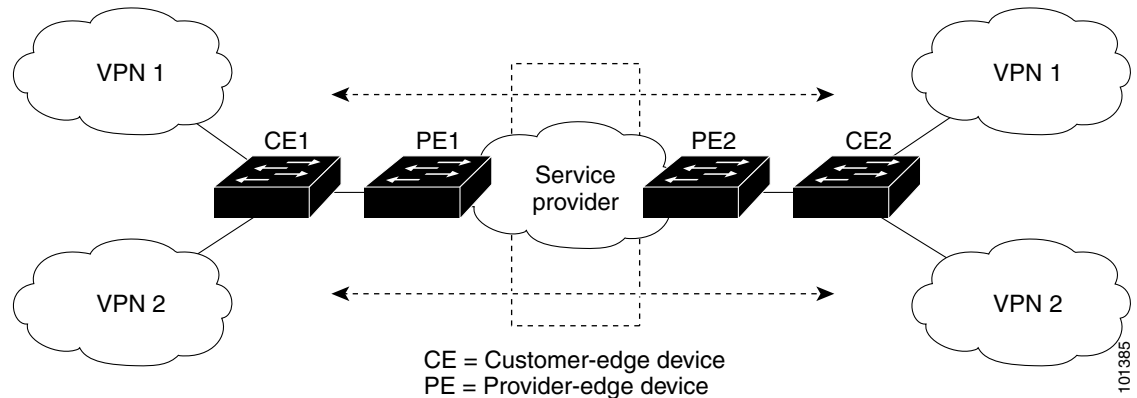
Multi-VRF CE interfaces must be Layer 3 interfaces.

Multi-VRF CE includes these devices:

- Customer edge (CE) devices provide customers access to the service-provider network over a data link to one or more provider edge routers. The CE device advertises the site's local routes to the router and learns the remote VPN routes from it. A Catalyst 3550 switch can be a CE.
- Provider edge (PE) routers exchange routing information with CE devices by using static routing or a routing protocol such as BGP, RIPv2, OSPF, or EIGRP. The PE is only required to maintain VPN routes for those VPNs to which it is directly attached, eliminating the need for the PE to maintain all of the service-provider VPN routes. Each PE router maintains a VRF for each of its directly connected sites. Multiple interfaces on a PE router can be associated with a single VRF if all of these sites participate in the same VPN. Each VPN is mapped to a specified VRF. After learning local VPN routes from CEs, a PE router exchanges VPN routing information with other PE routers by using internal BGP (IBPG).
- Provider routers or core routers are any routers in the service provider network that do not attach to CE devices.

With multi-VRF CE, multiple customers can share one CE, and only one physical link is used between the CE and the PE. The shared CE maintains separate VRF tables for each customer and switches or routes packets for each customer based on its own routing table. Multi-VRF CE extends limited PE functionality to a CE device, giving it the ability to maintain separate VRF tables to extend the privacy and security of a VPN to the branch office.

[Figure 32-5](#) shows a configuration using Catalyst 3550 switches as multiple virtual CEs. This scenario is suited for customers who have low bandwidth requirements for their VPN service, for example, small companies. In this case, multi-VRF CE support is required in the Catalyst 3550 switches. Because multi-VRF CE is a Layer 3 feature, each interface in a VRF must be a Layer 3 interface.

Figure 32-5 Catalyst 3550 Switches Acting as Multiple Virtual CEs

When the CE switch receives a command to add a Layer 3 interface to a VRF, it sets up the appropriate mapping between the VLAN ID and the policy label (PL) in multi-VRF-CE-related data structures and adds the VLAN ID and PL to the VLAN database.

When multi-VRF CE is configured, the Layer 3 forwarding table is conceptually partitioned into two sections:

- The multi-VRF CE routing section contains the routes from different VPNs.
- The global routing section contains routes to non-VPN networks, such as the Internet.

VLAN IDs from different VRFs are mapped into different policy labels, which are used to distinguish the VRFs during processing. If no route is found in the multi-VRF CE section of the Layer 3 forwarding table, the global routing section is used to determine the forwarding path. For each new VPN route learned, the Layer 3 setup function retrieves the policy label by using the VLAN ID of the ingress port and inserts the policy label and new route to the multi-VRF CE routing section. If the packet is received from a routed port, the port internal VLAN ID number is used; if the packet is received from an SVI, the VLAN number is used.

This is the packet-forwarding process in a multi-VRF-CE-enabled network:

- When the switch receives a packet from a VPN, the switch looks up the routing table based on the input policy label number. When a route is found, the switch forwards the packet to the PE.
- When the ingress PE receives a packet from the CE, it performs a VRF lookup. When a route is found, the router adds a corresponding MPLS label to the packet and sends it to the MPLS network.
- When an egress PE receives a packet from the network, it strips the label and uses the label to identify the correct VPN routing table. Then it performs the normal route lookup. When a route is found, it forwards the packet to the correct adjacency.
- When a CE receives a packet from an egress PE, it uses the input policy label to look up the correct VPN routing table. If a route is found, it forwards the packet within the VPN.

To configure VRF, you create a VRF table and specify the Layer 3 interface associated with the VRF. Then configure the routing protocols in the VPN and between the CE and the PE. BGP is the preferred routing protocol used to distribute VPN routing information across the provider's backbone.

The multi-VRF CE network has three major components:

- VPN route target communities—lists of all other members of a VPN community. You need to configure VPN route targets for each VPN community member.
- Multiprotocol BGP peering of VPN community PE routers—propagates VRF reachability information to all members of a VPN community. You need to configure BGP peering in all PE routers within a VPN community.
- VPN forwarding—transports all traffic between all VPN community members across a VPN service-provider network.

Default Multi-VRF CE Configuration

Table 32-13 shows the default VRF configuration.

Table 32-13 **Default VRF Configuration**

Feature	Default Setting
VRF	Disabled. No VRFs are defined.
Maps	No import maps, export maps, or route maps are defined.
VRF maximum routes	Fast Ethernet switches: 8000 Gigabit Ethernet switches: 12000.
Forwarding table	The default for an interface is the global routing table.

Multi-VRF CE Configuration Guidelines



Note

To use multi-VRF CE, you must have the enhanced multilayer software image installed on your switch.

These are considerations when configuring VRF in your network:

- A switch with multi-VRF CE is shared by multiple customers, and each customer has its own routing table.
- Because customers use different VRF tables, the same IP addresses can be reused. Overlapped IP addresses are allowed in different VPNs.
- Multi-VRF CE lets multiple customers share the same physical link between the PE and the CE. Trunk ports with multiple VLANs separate packets among customers. Each customer has its own VLAN.
- Multi-VRF CE does not support all MPLS-VRF functionality. It does not support label exchange, LDP adjacency, or labeled packets.
- For the PE router, there is no difference between using multi-VRF CE or using multiple CEs. In [Figure 32-5](#), multiple virtual Layer 3 interfaces are connected to the multi-VRF CE device.
- The switch supports configuring VRF by using physical ports, VLAN SVIs, or a combination of both. The SVIs can be connected through an access port or a trunk port.
- A customer can use multiple VLANs as long as they do not overlap with those of other customers. A customer's VLANs are mapped to a specific routing table ID that is used to identify the appropriate routing tables stored on the switch.
- To support multi-VRF CE, multiple routing tables are entered into the Layer 3 TCAM table. Because an extra field is needed in the routing table to identify the table to which a route belongs, you must modify the SDM template to enable the switch to support 144-bit Layer 3 TCAM. Use the **sdm prefer extended-match**, **sdm prefer access extended-match**, or **sdm prefer routing extended-match** global configuration command to reformat the TCAM space allocated to unicast routing in the default, access, or routing template, respectively. Reformatting the unicast routing TCAM halves the number of supported unicast routes in the template.



Note

For more information on the SDM templates, see the [“Optimizing System Resources for User-Selected Features”](#) section on page 6-26.

- A Catalyst 3550 switch supports one global network and up to seven VRFs. The total number of routes supported are limited by the size of the TCAM and specified in the SDM template.
- Most routing protocols (BGP, OSPF, RIP, and static routing) can be used between the CE and the PE. However, we recommend using external BGP (EBGP) for these reasons:
 - BGP does not require multiple algorithms to communicate with multiple CEs.
 - BGP is designed for passing routing information between systems run by different administrations.
 - BGP makes it easy to pass attributes of the routes to the CE.
- Multi-VRF-CE does not support EIGRP.
- Multi-VRF CE does not affect the packet switching rate.
- VPN multicast is not supported.

- You cannot configure the Web Cache Communication Protocol (WCCP) and multi-VRF CE on the same switch at the same time.
- When multi-VRF CE is configured, you cannot assign the same Hot Standby Routing Protocol (HSRP) standby address to two different VPNs.
- VRF and policy-based routing (PBR) are mutually-exclusive on a switch interface. You cannot enable VRF when PBR is enabled on an interface. In contrast, you cannot enable PBR when VRF is enabled on an interface.

Configuring VRFs

Beginning in privileged EXEC mode, follow these steps to configure one or more VRFs. For complete syntax and usage information for the commands, see the switch command reference for this release and the *Cisco IOS Switching Services Command Reference, Release 12.2*.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip routing	Enable IP routing.
Step 3	ip vrf <i>vrf-name</i>	Name the VRF, and enter VRF configuration mode.
Step 4	rd <i>route-distinguisher</i>	Create a VRF table by specifying a route distinguisher. Enter either an AS number and an arbitrary number (xxx:y) or an IP address and arbitrary number (A.B.C.D:y).
Step 5	route-target { export import both } <i>route-target-ext-community</i>	Create a list of import, export, or import and export route target communities for the specified VRF. Enter either an AS system number and an arbitrary number (xxx:y) or an IP address and an arbitrary number (A.B.C.D:y). The <i>route-target-ext-community</i> should be the same as the <i>route-distinguisher</i> entered in Step 4.
Step 6	import map <i>route-map</i>	(Optional) Associate a route map with the VRF.
Step 7	interface <i>interface-id</i>	Enter interface configuration mode and specify the Layer 3 interface to be associated with the VRF. The interface can be a routed port or SVI.
Step 8	ip vrf forwarding <i>vrf-name</i>	Associate the VRF with the Layer 3 interface.
Step 9	end	Return to privileged EXEC mode.
Step 10	show ip vrf [brief detail interfaces] <i>[vrf-name]</i>	Verify the configuration. Display information about the configured VRFs.
Step 11	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no ip vrf** *vrf-name* global configuration command to delete a VRF and to remove all interfaces from it. Use the **no ip vrf forwarding** interface configuration command to remove an interface from the VRF.

Configuring a VPN Routing Session

Routing within the VPN can be configured with any supported routing protocol (RIP, OSPF, EIGRP, or BGP) or with static routing. The configuration shown here is for OSPF, but the process is the same for other protocols.

Beginning in privileged EXEC mode, follow these steps to configure OSPF in the VPN:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router ospf <i>process-id</i> vrf <i>vrf-name</i>	Enable OSPF routing, specify a VPN forwarding table, and enter router configuration mode.
Step 3	log-adjacency-changes	(Optional) Log changes in the adjacency state. This is the default state.
Step 4	redistribute bgp <i>autonomous-system-number</i> subnets	Set the switch to redistribute information from the BGP network to the OSPF network.
Step 5	network <i>network-number</i> area <i>area-id</i>	Define a network address and mask on which OSPF runs and the area ID for that network address.
Step 6	end	Return to privileged EXEC mode.
Step 7	show ip ospf <i>process-id</i>	Verify the configuration of the OSPF network.
Step 8	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no router ospf** *process-id* **vrf** *vrf-name* global configuration command to disassociate the VPN forwarding table from the OSPF routing process.

Configuring BGP PE to CE Routing Sessions

Beginning in privileged EXEC mode, follow these steps to configure a BGP PE to CE routing session:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router bgp <i>autonomous-system-number</i>	Configure the BGP routing process with the AS number passed to other BGP routers, and enter router configuration mode.
Step 3	network <i>network-number</i> mask <i>network-mask</i>	Specify a network and mask to announce using BGP.
Step 4	redistribute ospf <i>process-id</i> match internal	Set the switch to redistribute OSPF internal routes.
Step 5	network <i>network-number</i> area <i>area-id</i>	Define a network address and mask on which OSPF runs and the area ID for that network address.
Step 6	address-family ipv4 vrf <i>vrf-name</i>	Define BGP parameters for PE to CE routing sessions, and enter VRF address-family mode.
Step 7	neighbor <i>address</i> remote-as <i>as-number</i>	Define a BGP session between PE and CE routers.
Step 8	neighbor <i>address</i> activate	Activate the advertisement of the IPv4 address family.

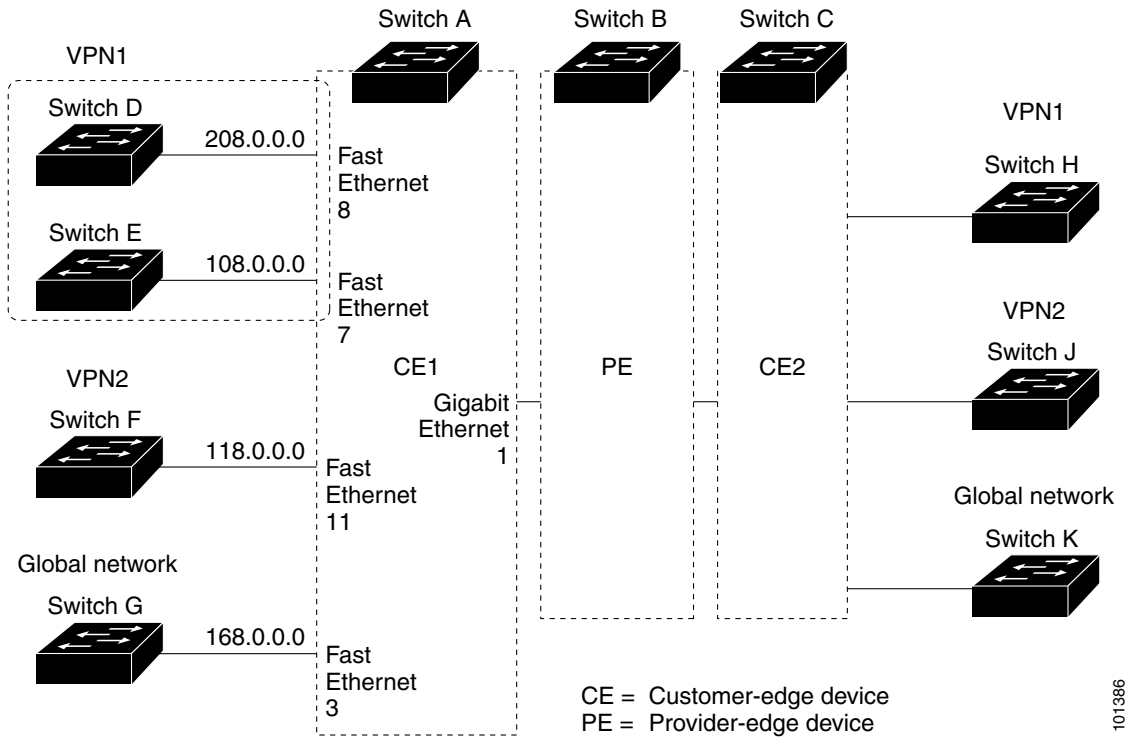
	Command	Purpose
Step 9	end	Return to privileged EXEC mode.
Step 10	show ip bgp [ipv4] [neighbors]	Verify BGP configuration.
Step 11	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no router bgp autonomous-system-number** global configuration command to delete the BGP routing process. Use the command with keywords to delete routing characteristics.

Multi-VRF CE Configuration Example

Figure 32-6 is a simplified example of the physical connections in a network similar to that in Figure 32-5. OSPF is the protocol used in VPN1, VPN2, and the global network. BGP is used in the CE to PE connections. The examples following the illustration show how to configure a Catalyst 3550 switch as CE Switch A, and the VRF configuration for customer switches D and F. Commands for configuring CE Switch C and the other customer switches are not included but would be similar. The example also includes commands for configuring traffic to Switch A for a Catalyst 6000 or Catalyst 6500 switch acting as a PE router.

Figure 32-6 Multi-VRF CE Configuration Example



Configuring Switch A

On Switch A, enable routing and configure VRF.

```
Switch# configure terminal
```

Enter configuration commands, one per line. End with CNTL/Z.

```
Switch(config)# ip routing
Switch(config)# ip vrf v11
Switch(config-vrf)# rd 800:1
Switch(config-vrf)# route-target export 800:1
Switch(config-vrf)# route-target import 800:1
Switch(config-vrf)# exit
Switch(config)# ip vrf v12
Switch(config-vrf)# rd 800:2
Switch(config-vrf)# route-target export 800:2
Switch(config-vrf)# route-target import 800:2
Switch(config-vrf)# exit
```

Configure the loopback and physical interfaces on Switch A. Gigabit Ethernet port 1 is a trunk connection to the PE. Fast Ethernet ports 8 and 11 connect to VPNs:

```
Switch(config)# interface loopback1
Switch(config-if)# ip vrf forwarding v11
Switch(config-if)# ip address 8.8.1.8 255.255.255.0
Switch(config-if)# exit

Switch(config)# interface loopback2
Switch(config-if)# ip vrf forwarding v12
Switch(config-if)# ip address 8.8.2.8 255.255.255.0
Switch(config-if)# exit

Switch(config)# interface gigabitethernet0/5
Switch(config-if)# switchport trunk encapsulation dot1q
Switch(config-if)# switchport mode trunk
Switch(config-if)# no ip address
Switch(config-if)# exit

Switch(config)# interface fastethernet0/8
Switch(config-if)# switchport access vlan 208
Switch(config-if)# no ip address
Switch(config-if)# exit

Switch(config)# interface fastethernet0/11
Switch(config-if)# switchport trunk encapsulation dot1q
Switch(config-if)# switchport mode trunk
Switch(config-if)# no ip address
Switch(config-if)# exit
```

Configure the VLANs used on Switch A. VLAN 10 is used by VRF 11 between the CE and the PE. VLAN 20 is used by VRF 12 between the CE and the PE. VLANs 118 and 208 are used for VRF for the VPNs that include Switch F and Switch D, respectively:

```
Switch(config)# interface vlan10
Switch(config-if)# ip vrf forwarding v11
Switch(config-if)# ip address 38.0.0.8 255.255.255.0
Switch(config-if)# exit

Switch(config)# interface vlan20
Switch(config-if)# ip vrf forwarding v12
Switch(config-if)# ip address 83.0.0.8 255.255.255.0
Switch(config-if)# exit

Switch(config)# interface vlan118
Switch(config-if)# ip vrf forwarding v12
Switch(config-if)# ip address 118.0.0.8 255.255.255.0
Switch(config-if)# exit

Switch(config)# interface vlan208
```

```
Switch(config-if)# ip vrf forwarding v11
Switch(config-if)# ip address 208.0.0.8 255.255.255.0
Switch(config-if)# exit
```

Configure OSPF routing in VPN1 and VPN2.

```
Switch(config)# router ospf 1 vrf v11
Switch(config-router)# redistribute bgp 800 subnets
Switch(config-router)# network 208.0.0.0 0.0.0.255 area 0
Switch(config-router)# exit
Switch(config)# router ospf 2 vrf v12
Switch(config-router)# redistribute bgp 800 subnets
Switch(config-router)# network 118.0.0.0 0.0.0.255 area 0
Switch(config-router)# exit
```

Configure BGP for CE to PE routing.

```
Switch(config)# router bgp 800
Switch(config-router)# address-family ipv4 vrf v12
Switch(config-router-af)# redistribute ospf 2 match internal
Switch(config-router-af)# neighbor 83.0.0.3 remote-as 100
Switch(config-router-af)# neighbor 83.0.0.3 activate
Switch(config-router-af)# network 8.8.2.0 mask 255.255.255.0
Switch(config-router-af)# exit

Switch(config-router)# address-family ipv4 vrf v11
Switch(config-router-af)# redistribute ospf 1 match internal
Switch(config-router-af)# neighbor 38.0.0.3 remote-as 100
Switch(config-router-af)# neighbor 38.0.0.3 activate
Switch(config-router-af)# network 8.8.1.0 mask 255.255.255.0
Switch(config-router-af)# end
```

Configuring Switch D

Switch D belongs to VPN 1 and is connected to Switch A by using these commands.

```
Switch# configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)# ip routing
Switch(config)# interface fastethernet0/2
Switch(config-if)# no switchport
Switch(config-if)# ip address 208.0.0.20 255.255.255.0
Switch(config-if)# exit

Switch(config)# router ospf 101
Switch(config-router)# network 208.0.0.0 0.0.0.255 area 0
Switch(config-router)# end
```

Configuring Switch F

Switch F belongs to VPN 2 and is connected to Switch A by using these commands.

```
Switch# configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)# ip routing
Switch(config)# interface fastethernet0/1
Switch(config-if)# switchport trunk encapsulation dot1q
Switch(config-if)# switchport mode trunk
Switch(config-if)# no ip address
Switch(config-if)# exit

Switch(config)# interface vlan118
Switch(config-if)# ip address 118.0.0.11 255.255.255.0
Switch(config-if)# exit
```

```
Switch(config)# router ospf 101
Switch(config-router)# network 118.0.0.0 0.0.0.255 area 0
Switch(config-router)# end
```

Configuring the PE Switch B

On Switch B (the PE router), these commands only configure the connections to the CE device, Switch A.

```
Router# configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Router(config)# ip vrf v1
Router(config-vrf)# rd 100:1
Router(config-vrf)# route-target export 100:1
Router(config-vrf)# route-target import 100:1
Router(config-vrf)# exit

Router(config)# ip vrf v2
Router(config-vrf)# rd 100:2
Router(config-vrf)# route-target export 100:2
Router(config-vrf)# route-target import 100:2
Router(config-vrf)# exit

Router(config)# ip cef
Router(config)# interface Loopback1
Router(config-if)# ip vrf forwarding v1
Router(config-if)# ip address 3.3.1.3 255.255.255.0
Router(config-if)# exit

Router(config)# interface Loopback2
Router(config-if)# ip vrf forwarding v2
Router(config-if)# ip address 3.3.2.3 255.255.255.0
Router(config-if)# exit

Router(config)# interface gigabitethernet1/0.10
Router(config-if)# encapsulation dot1q 10
Router(config-if)# ip vrf forwarding v1
Router(config-if)# ip address 38.0.0.3 255.255.255.0
Router(config-if)# exit

Router(config)# interface gigabitethernet1/0.20
Router(config-if)# encapsulation dot1q 20
Router(config-if)# ip vrf forwarding v2
Router(config-if)# ip address 83.0.0.3 255.255.255.0
Router(config-if)# exit

Router(config)# router bgp 100
Router(config-router)# address-family ipv4 vrf v2
Router(config-router-af)# neighbor 83.0.0.8 remote-as 800
Router(config-router-af)# neighbor 83.0.0.8 activate
Router(config-router-af)# network 3.3.2.0 mask 255.255.255.0
Router(config-router-af)# exit
Router(config-router)# address-family ipv4 vrf v1
Router(config-router-af)# neighbor 83.0.0.8 remote-as 800
Router(config-router-af)# neighbor 83.0.0.8 activate
Router(config-router-af)# network 3.3.1.0 mask 255.255.255.0
Router(config-router-af)# end
```

Displaying Multi-VRF CE Status

You can use the privileged EXEC commands in [Table 32-14](#) to display information about multi-VRF CE configuration and status.

Table 32-14 Commands for Displaying Multi-VRF CE Information

Command	Purpose
<code>show ip protocols vrf vrf-name</code>	Display routing protocol information associated with a VRF.
<code>show ip route vrf vrf-name [connected] [protocol [as-number]] [list] [mobile] [odr] [profile] [static] [summary] [supernets-only]</code>	Display IP routing table information associated with a VRF.
<code>show ip vrf [brief detail interfaces] [vrf-name]</code>	Display information about the defined VRF instances.

For more information about the information in the displays, see the *Cisco IOS Switching Services Command Reference, Release 12.2*.

Configuring Protocol-Independent Features

This section describes how to configure IP routing protocol-independent features. These features are available on switches running the IP base image or the IP services image, but protocol-related features with the IP base image are available only for RIP. For a complete description of the IP routing protocol-independent commands in this chapter, see the “IP Routing Protocol-Independent Commands” chapter of the *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*.

This section includes these procedures:

- [Configuring Cisco Express Forwarding, page 32-72](#)
- [Configuring the Number of Equal-Cost Routing Paths, page 32-74](#)
- [Configuring Static Unicast Routes, page 32-74](#)
- [Specifying Default Routes and Networks, page 32-75](#)
- [Using Route Maps to Redistribute Routing Information, page 32-76](#)
- [Configuring Policy-Based Routing, page 32-79](#)
- [Filtering Routing Information, page 32-82](#)
- [Managing Authentication Keys, page 32-84](#)

Configuring Cisco Express Forwarding

Cisco Express Forwarding (CEF) is a Layer 3 IP switching technology used to optimize network performance. CEF implements an advanced IP look-up and forwarding algorithm to deliver maximum Layer 3 switching performance. CEF is less CPU-intensive than fast switching route caching, allowing more CPU processing power to be dedicated to packet forwarding. In dynamic networks, fast switching cache entries are frequently invalidated because of routing changes, which can cause traffic to be

process-switched by using the routing table, instead of fast-switched by using the route cache. CEF uses the forwarding information base (FIB) lookup table to perform destination-based switching of IP packets.

The two main components in CEF are the FIB and adjacency tables.

- The FIB is similar to a routing table or information base and maintains a mirror image of the forwarding information in the IP routing table. When routing or topology changes occur in the network, the IP routing table is updated, and those changes are reflected in the FIB. The FIB maintains next-hop address information based on the information in the IP routing table. Because the FIB contains all known routes that exist in the routing table, CEF eliminates route cache maintenance, is more efficient for switching traffic, and is not affected by traffic patterns.
- Nodes in the network are said to be adjacent if they can reach each other with a single hop across a link layer. CEF uses adjacency tables to prepend Layer 2 addressing information. The adjacency table maintains Layer 2 next-hop addresses for all FIB entries.

Because the switch uses an Application Specific Integrated Circuit (ASIC) to achieve Gigabit-speed line rate IP traffic, CEF forwarding applies only to the software forwarding path, that is, traffic that is forwarded by the CPU.

CEF is enabled globally by default. If for some reason it is disabled, you can re-enable it by using the **ip cef** global configuration command.

The default configuration, which we recommend, is CEF enabled on all Layer 3 interfaces. On the switch, you can use the **no ip route-cache cef** interface configuration command to disable CEF for traffic that is being forwarded by software. This command does not affect the hardware forwarding path. Disabling CEF and using the **debug ip packet detail** privileged EXEC command can be useful when you want to debug software-forwarded traffic. You can enable CEF on an interface for the software-forwarding path by using the **ip route-cache cef** interface configuration command.

Beginning in privileged EXEC mode, follow these steps to enable CEF on an interface for software-forwarded traffic after it has been disabled:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Enter interface configuration mode, and specify the Layer 3 interface to configure.
Step 3	ip route-cache cef	Enable CEF on the interface for software-forwarded traffic.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip cef	Display the CEF status on all interfaces.
Step 6	show adjacency	Display CEF adjacency table information.
Step 7	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To disable CEF on an interface for software-forwarded traffic, use the **no ip route-cache cef** interface configuration command.

Configuring the Number of Equal-Cost Routing Paths

When a router has two or more routes to the same network with the same metrics, these routes can be thought of as having an equal cost. The term *parallel path* is another way to refer to occurrences of equal-cost routes in a routing table. If a router has two or more equal-cost paths to a network, it can use them concurrently. Parallel paths provide redundancy in case of a circuit failure and also enable a router to load balance packets over the available paths for more efficient use of available bandwidth.

Even though the router automatically learns about and configures equal-cost routes, you can control the maximum number of parallel paths supported by an IP routing protocol in its routing table. Although the switch software allows a maximum of 32 equal-cost routes, the switch hardware will never use more than 8 paths per route.

Beginning in privileged EXEC mode, follow these steps to change the maximum number of parallel paths installed in a routing table from the default:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router {bgp rip ospf eigrp}	Enter router configuration mode.
Step 3	maximum-paths <i>maximum</i>	Set the maximum number of parallel paths for the protocol routing table. The range is from 1 to 8; the default is 4 for most IP routing protocols, but only 1 for BGP.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip protocols	Verify the setting in the <i>Maximum path</i> field.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no maximum-paths** router configuration command to restore the default value.

Configuring Static Unicast Routes

Static unicast routes are user-defined routes that cause packets moving between a source and a destination to take a specified path. Static routes can be important if the router cannot build a route to a particular destination and are useful for specifying a gateway of last resort to which all unroutable packets are sent.

Beginning in privileged EXEC mode, follow these steps to configure a static unicast route:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip route <i>prefix mask</i> {<i>address</i> <i>interface</i>} [<i>distance</i>]	Establish a static route.
Step 3	end	Return to privileged EXEC mode.
Step 4	show ip route	Display the current state of the routing table to verify the configuration.
Step 5	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no ip route *prefix mask*** global configuration command to remove a static route.

The switch retains static routes until you remove them. However, you can override static routes with dynamic routing information by assigning administrative distance values. Each dynamic routing protocol has a default administrative distance, as listed in [Table 32-15](#). If you want a static route to be overridden by information from a dynamic routing protocol, set the administrative distance of the static route higher than that of the dynamic protocol.

Table 32-15 *Dynamic Routing Protocol Default Administrative Distances*

Route Source	Default Distance
Connected interface	0
Static route	1
EIGRP summary route	5
External BGP	20
Internal Enhanced IGRP	90
OSPF	110
RIP	120
EIGRP summary route	170
Internal BGP	200
Unknown	225

Static routes that point to an interface are advertised through RIP and other dynamic routing protocols, whether or not static **redistribute** router configuration commands were specified for those routing protocols. These static routes are advertised because static routes that point to an interface are considered in the routing table to be connected and hence lose their static nature. However, if you define a static route to an interface that is not one of the networks defined in a network command, no dynamic routing protocols advertise the route unless a **redistribute** static command is specified for these protocols.

When an interface goes down, all static routes through that interface are removed from the IP routing table. When the software can no longer find a valid next hop for the address specified as the forwarding router's address in a static route, the static route is also removed from the IP routing table.

Specifying Default Routes and Networks

A router might not be able to determine the routes to all other networks. To provide complete routing capability, you can use some routers as smart routers and give the remaining routers default routes to the smart router. (Smart routers have routing table information for the entire internetwork.) These default routes can be dynamically learned or can be configured in the individual routers. Most dynamic interior routing protocols include a mechanism for causing a smart router to generate dynamic default information that is then forwarded to other routers.

If a router has a directly connected interface to the specified default network, the dynamic routing protocols running on that device generate a default route. In RIP, it advertises the pseudonetwork 0.0.0.0.

A router that is generating the default for a network also might need a default of its own. One way a router can generate its own default is to specify a static route to the network 0.0.0.0 through the appropriate device.

Beginning in privileged EXEC mode, follow these steps to define a static route to a network as the static default route:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip default-network <i>network number</i>	Specify a default network.
Step 3	end	Return to privileged EXEC mode.
Step 4	show ip route	Display the selected default route in the gateway of last resort display.
Step 5	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no ip default-network** *network number* global configuration command to remove the route.

When default information is passed through a dynamic routing protocol, no further configuration is required. The system periodically scans its routing table to choose the optimal default network as its default route. Cisco routers use administrative distance and metric information to determine the default route or the gateway of last resort.

If dynamic default information is not being passed to the system, candidates for the default route are specified with the **ip default-network** global configuration command. If this network appears in the routing table from any source, it is flagged as a possible choice for the default route. If the router has no interface on the default network, but does have a path to it, the network is considered as a possible candidate, and the gateway to the best default path becomes the gateway of last resort.

Using Route Maps to Redistribute Routing Information

The switch can run multiple routing protocols simultaneously, and it can redistribute information from one routing protocol to another. Redistributing information from one routing protocol to another applies to all supported IP-based routing protocols.

You can conditionally control the redistribution of routes between routing domains by defining route maps between the two domains. The **match** and **set** route-map configuration commands define the condition portion of a route map. The **match** command specifies that a criterion must be matched; the **set** command specifies an action to be taken if the routing update meets the conditions defined by the match command. Although redistribution is a protocol-independent feature, some of the **match** and **set** route-map configuration commands are specific to a particular protocol.

One or more **match** commands and one or more **set** commands follow a **route-map** command. If there are no **match** commands, everything matches. If there are no **set** commands, nothing is done, other than the match. Therefore, you need at least one **match** or **set** command.



Note

Although each of Steps 3 through 16 in the following section is optional, you must enter at least one **match** route-map configuration command and one **set** route-map configuration command. For complete syntax information for the command, see the *Cisco IOS IP and IP Routing Command Reference, Release 12.2*.

Beginning in privileged EXEC mode, follow these steps to configure a route map for redistribution:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	route-map <i>map-tag</i> [permit deny] [<i>sequence number</i>]	Define any route maps used to control redistribution and enter route-map configuration mode. <i>map-tag</i> —A meaningful name for the route map. The redistribute router configuration command uses this name to reference the route map. Multiple route maps might share the same map tag name. (Optional) If permit is specified and the match criteria are met for this route map, the route is redistributed as controlled by the set actions. If deny is specified, the route is not redistributed. <i>sequence number</i> (Optional)— Number that indicates the position of a new route map in the list of route maps already configured with the same name.
Step 3	match as-path <i>path-list-number</i>	Match a BGP AS path access list.
Step 4	match community-list <i>community-list-number</i> [exact]	Match a BGP community list.
Step 5	match ip address { <i>access-list-number</i> <i>access-list-name</i> } [... <i>access-list-number</i> ... <i>access-list-name</i>]	Match a standard access list by specifying the name or number. It can be an integer from 1 to 199.
Step 6	match ip next-hop { <i>access-list-number</i> <i>access-list-name</i> } [... <i>access-list-number</i> ... <i>access-list-name</i>]	Match a next-hop router address passed by one of the access lists specified (numbered from 1 to 199).
Step 7	match tag <i>tag value</i> [... <i>tag-value</i>]	Match the specified tag value in a list of one or more route tag values. Each can be an integer from 0 to 4294967295.
Step 8	match interface <i>type number</i> [... <i>type number</i>]	Match the specified next hop route out one of the specified interfaces.
Step 9	match ip route-source { <i>access-list-number</i> <i>access-list-name</i> } [... <i>access-list-number</i> ... <i>access-list-name</i>]	Match the address specified by the specified advertised access lists.
Step 10	match route-type { local internal external [type-1 type-2]	Match the specified route-type : • local—Locally generated BGP routes. • internal—OSPF intra-area and interarea routes or EIGRP internal routes. • external—OSPF external routes (Type 1 or Type 2) or EIGRP external routes.
Step 11	set dampening <i>half-life</i> <i>reuse</i> <i>suppress</i> <i>max-suppress-time</i>	Set BGP route dampening factors.
Step 12	set local-preference <i>value</i>	Assign a value to a local BGP path.
Step 13	set origin { igp egp as incomplete }	Set the BGP origin code.
Step 14	set as-path { tag prepend <i>as-path-string</i> }	Modify the BGP autonomous system path.

	Command	Purpose
Step 15	set level { level-1 level-2 level-1-2 stub-area backbone }	Set the level for routes that are advertised into the specified area of the routing domain. The stub-area and backbone are OSPF NSSA and backbone areas.
Step 16	set metric <i>metric value</i>	Set the metric value to give the redistributed routes (for any protocol except EIGRP). The <i>metric value</i> is an integer from –294967295 to 294967295.
Step 17	set metric <i>bandwidth delay reliability loading mtu</i>	Set the metric value to give the redistributed routes (for EIGRP only): <i>bandwidth</i>—Metric value or in kilobits per second in the range 0 to 4294967295. <i>delay</i>—Route delay in tens of microseconds in the range 0 to 4294967295. <i>reliability</i>—Likelihood of successful packet transmission expressed as a number between 0 (no reliability) and 255 (100 percent reliability). <i>loading</i>—Effective bandwidth of the route expressed as a number from 0 to 255 (100 percent loading). <i>mtu</i>—Maximum transmission unit (MTU) size of the route in bytes in the range 0 to 4294967295.
Step 18	set metric-type { internal external type-1 type-2 }	Set the metric type to give redistributed routes.
Step 19	set metric-type internal	Set the multi-exit discriminator (MED) value on prefixes advertised to External BGP neighbor to match the IGP metric of the next hop
Step 20	set weight	Set the BGP weight for the routing table. The value can be from 1 to 65535.
Step 21	end	Return to privileged EXEC mode.
Step 22	show route-map	Display all route maps configured or only the one specified to verify configuration.
Step 23	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To delete an entry, use the **no route-map** *map tag* global configuration command or the **no match** or **no set** route-map configuration commands.

You can distribute routes from one routing domain into another and control route distribution.

Beginning in privileged EXEC mode, follow these steps to control route redistribution. Note that the keywords are the same as defined in the previous procedure.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router { bgp rip ospf eigrp }	Enter router configuration mode.
Step 3	redistribute <i>protocol</i> [<i>process-id</i>] { level-1 level-1-2 level-2 } [metric <i>metric-value</i>] [metric-type <i>type-value</i>] [match internal external <i>type-value</i>] [tag <i>tag-value</i>] [route-map <i>map-tag</i>] [weight <i>weight</i>] [subnets]	Redistribute routes from one routing protocol to another routing protocol.

	Command	Purpose
Step 4	default-metric <i>number</i>	Cause the current routing protocol to use the same metric value for all redistributed routes (BGP, RIP, and OSPF).
Step 5	default-metric <i>bandwidth delay reliability loading mtu</i>	Cause the EIGRP routing protocol to use the same metric value for all non-EIGRP redistributed routes.
Step 6	end	Return to privileged EXEC mode.
Step 7	show route-map	Display all route maps configured or only the one specified to verify configuration.
Step 8	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To disable redistribution, use the **no** form of the commands.

The metrics of one routing protocol do not necessarily translate into the metrics of another. For example, the RIP metric is a hop count. In these situations, an artificial metric is assigned to the redistributed route. Uncontrolled exchanging of routing information between different routing protocols can create routing loops and seriously degrade network operation.

If you have not defined a default redistribution metric that replaces metric conversion, some automatic metric translations occur between routing protocols:

- RIP can automatically redistribute static routes. It assigns static routes a metric of 1 (directly connected).
- Any protocol can redistribute other routing protocols if a default mode is in effect.

Configuring Policy-Based Routing

You can use policy-based routing (PBR) to configure a defined policy for traffic flows. By using PBR, you can have more control over routing by reducing the reliance on routes derived from routing protocols. PBR can determine and implement routing policies that allow or deny paths based on:

- Identity of a particular end system
- Application
- Protocol

You can use PBR to provide equal-access and source-sensitive routing, routing based on interactive versus batch traffic, or routing based on dedicated links. For example, you could transfer stock records to a corporate office on a high-bandwidth, high-cost link for a short time while transmitting routine application data such as e-mail over a low-bandwidth, low-cost link.

PBR is applied to incoming packets. All packets received on an interface with PBR enabled are considered for PBR. The switch passes the packets through route maps. Based on the criteria defined in the route maps, packets are forwarded (routed) to the appropriate next hop. For more information about configuring route maps see the [“Using Route Maps to Redistribute Routing Information”](#) section on page 32-76.

For details about PBR commands and keywords, see the *Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*. For a list of PBR commands not supported by the switch, see [Appendix C, “Unsupported CLI Commands in Cisco IOS Release 12.2\(25\)SEE.”](#)

PBR Configuration Guidelines

Before configuring PBR, you should be aware of this information:

- To use PBR, you must have the IP services image installed on your switch.
- Multicast traffic is not policy-routed. PBR applies to only to unicast traffic.
- You can enable PBR on a routed port, an SVI, or an EtherChannel port channel in Layer 3 mode.
- You can define a maximum of 247 IP policy route-maps on the switch.
- VRF and PBR are mutually-exclusive on a switch interface. You cannot enable VRF when PBR is enabled on an interface. In contrast, you cannot enable PBR when VRF is enabled on an interface.
- WCCP and PBR are mutually-exclusive on a switch interface. You cannot enable WCCP when PBR is enabled on an interface. In contrast, you cannot enable PBR when WCCP is enabled on an interface.
- The number of TCAM entries used by PBR depends on the route-map itself, the ACLs used, and the order of the ACLs and route-map entries.
- You must modify the SDM template to enable the switch to support the 144-bit Layer 3 TCAM. Use the **sdm prefer extended-match**, **sdm prefer access extended-match**, or the **sdm prefer routing extended-match** global configuration commands to reformat the TCAM space allocated to unicast routing in the default, access, or routing template, respectively. Reformatting the unicast routing TCAM reduces by half the number of supported unicast routes in the template.

See the [“Optimizing System Resources for User-Selected Features” section on page 6-26](#) and the [“Displaying ACL Resource Usage and Configuration Problems” section on page 29-43](#) for more information about managing the memory resources in the switch.

Enabling PBR

By default, PBR is disabled on the switch. To enable PBR, you must create a route map that specifies the match criteria and the resulting action if all of the match clauses are met. Then, you must enable PBR for that route map on an interface. All packets arriving on the specified interface matching the match clauses will be subject to PBR.

Beginning in privileged EXEC mode, follow these steps to configure PBR:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	route-map <i>map-tag</i> [permit deny] [<i>sequence number</i>]	Define any route maps used to control where packets are output and enter route-map configuration mode. <i>map-tag</i> —A meaningful name for the route map. The ip policy route-map interface configuration command uses this name to reference the route map. Multiple route maps might share the same map tag name. (Optional) If permit is specified and the match criteria are met for this route map, the route is policy-routed as controlled by the set actions. If deny is specified, the route is not policy-routed. <i>sequence number</i> (Optional)— Number that shows the position of a new route map in the list of route maps already configured with the same name.
Step 3	match ip address { <i>access-list-number</i> <i>access-list-name</i> } [... <i>access-list-number</i> ... <i>access-list-name</i>]	Match the source and destination IP address that is permitted by one or more standard or extended access lists. If you do not specify a match command, the route map applies to all packets.
Step 4	set ip next-hop <i>ip-address</i> [... <i>ip-address</i>]	Specify the action to take on the packets that match the criteria. Set next hop to which to route the packet (the next hop must be adjacent).
Step 5	interface <i>interface-id</i>	Enter interface configuration mode, and specify the interface to configure.
Step 6	ip policy route-map <i>map-tag</i>	Enable PBR on a Layer 3 interface, and identify the route map to use. You can configure only one route map on an interface. However, you can have multiple route map entries with different sequence numbers. These entries are evaluated in sequence number order until the first match. If there is no match, packets are routed as usual.
Step 7	ip route-cache policy	(Optional) Enable fast-switching PBR. You must first enable PBR before enabling fast-switching PBR.
Step 8	exit	Return to global configuration mode.
Step 9	ip local policy route-map <i>map-tag</i>	(Optional) Enable local PBR to perform policy-based routing on packets originating at the switch. This applies to packets generated by the switch and not to incoming packets.

	Command	Purpose
Step 10	end	Return to privileged EXEC mode.
Step 11	show route-map <i>[map-name]</i>	Display all route maps configured or only the one specified to verify configuration.
Step 12	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no route-map** *map-tag* global configuration command or the **no match** or **no set** route-map configuration commands to delete an entry. Use the **no ip policy route-map** *map-tag* interface configuration command to disable PBR on an interface.

Filtering Routing Information

You can filter routing protocol information by performing the tasks described in this section.



Note

When routes are redistributed between OSPF processes, no OSPF metrics are preserved.

Setting Passive Interfaces

To prevent other routers on a local network from dynamically learning about routes, you can use the **passive-interface** router configuration command to keep routing update messages from being sent through a router interface. When you use this command in the OSPF protocol, the interface address you specify as passive appears as a stub network in the OSPF domain. OSPF routing information is neither sent nor received through the specified router interface.

In networks with many interfaces, to avoid having to manually set them as passive, you can set all interfaces to be passive by default by using the **passive-interface default** router configuration command and manually setting interfaces where adjacencies are desired.

Beginning in privileged EXEC mode, follow these steps to configure passive interfaces:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router { bgp rip ospf eigrp }	Enter router configuration mode.
Step 3	passive-interface <i>interface-id</i>	Suppress sending routing updates through the specified Layer 3 interface.
Step 4	passive-interface default	(Optional) Set all interfaces as passive by default.
Step 5	no passive-interface <i>interface type</i>	(Optional) Activate only those interfaces that need to have adjacencies sent.
Step 6	network <i>network-address</i>	(Optional) Specify the list of networks for the routing process. The <i>network-address</i> is an IP address.
Step 7	end	Return to privileged EXEC mode.
Step 8	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use a network monitoring privileged EXEC command such as **show ip ospf interface** to verify the interfaces that you enabled as passive, or use the **show ip interface** privileged EXEC command to verify the interfaces that you enabled as active.

To re-enable the sending of routing updates, use the **no passive-interface** *interface-id* router configuration command. The **default** keyword sets all interfaces as passive by default. You can then configure individual interfaces where you want adjacencies by using the **no passive-interface** router configuration command. The **default** keyword is useful in Internet service provider and large enterprise networks where many of the distribution routers have more than 200 interfaces.

Controlling Advertising and Processing in Routing Updates

You can use the **distribute-list** router configuration command with access control lists to suppress routes from being advertised in routing updates and to prevent other routers from learning one or more routes. When used in OSPF, this feature applies to only external routes, and you cannot specify an interface name.

You can also use a **distribute-list** router configuration command to avoid processing certain routes listed in incoming updates. (This feature does not apply to OSPF.)

Beginning in privileged EXEC mode, follow these steps to control the advertising or processing of routing updates:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router {bgp rip ospf eigrp}	Enter router configuration mode.
Step 3	distribute-list { <i>access-list-number</i> <i>access-list-name</i> } out [<i>interface-name</i> <i>routing process</i> <i>autonomous-system-number</i>]	Permit or deny routes from being advertised in routing updates, depending upon the action listed in the access list.
Step 4	distribute-list { <i>access-list-number</i> <i>access-list-name</i> } in [<i>type-number</i>]	Suppress processing in routes listed in updates.
Step 5	end	Return to privileged EXEC mode.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

Use the **no distribute-list in** router configuration command to change or cancel a filter. To cancel suppression of network advertisements in updates, use the **no distribute-list out** router configuration command.

Filtering Sources of Routing Information

Because some routing information might be more accurate than others, you can use filtering to prioritize information coming from different sources. An *administrative distance* is a rating of the trustworthiness of a routing information source, such as a router or group of routers. In a large network, some routing protocols can be more reliable than others. By specifying administrative distance values, you enable the router to intelligently discriminate between sources of routing information. The router always picks the route whose routing protocol has the lowest administrative distance. [Table 32-15 on page 32-75](#) shows the default administrative distances for various routing information sources.

Because each network has its own requirements, there are no general guidelines for assigning administrative distances.

Beginning in privileged EXEC mode, follow these steps to filter sources of routing information:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	router {bgp rip ospf eigrp}	Enter router configuration mode.
Step 3	distance weight {ip-address {ip-address mask}} [ip access list]	Define an administrative distance. <i>weight</i> —Administrative distance as an integer from 10 to 255. Used alone, <i>weight</i> specifies a default administrative distance used when no other specification exists for a routing information source. Routes with a distance of 255 are not installed in the routing table. (Optional) <i>ip access list</i> —IP standard or extended access list to be applied to incoming routing updates.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip protocols	Display the default administrative distance for a specified routing process.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To remove a distance definition, use the **no distance** router configuration command.

Managing Authentication Keys

Key management is a method of controlling authentication keys used by routing protocols. Not all protocols can use key management. Authentication keys are available for EIGRP and RIP Version 2.

Before you manage authentication keys, you must enable authentication. See the appropriate protocol section to see how to enable authentication for that protocol. To manage authentication keys, define a key chain, identify the keys that belong to the key chain, and specify how long each key is valid. Each key has its own key identifier (specified with the **key number** key chain configuration command), which is stored locally. The combination of the key identifier and the interface associated with the message uniquely identifies the authentication algorithm and Message Digest 5 (MD5) authentication key in use.

You can configure multiple keys with lifetimes. Only one authentication packet is sent, regardless of how many valid keys exist. The software examines the key numbers in order from lowest to highest and uses the first valid key it encounters. The lifetimes allow for overlap during key changes. Note that the router must know these lifetimes.

Beginning in privileged EXEC mode, follow these steps to manage authentication keys:

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	key chain name-of-chain	Identify a key chain, and enter key chain configuration mode.
Step 3	key number	Identify the key number. The range is 0 to 2147483647.
Step 4	key-string text	Identify the key string. The string can contain from 1 to 80 uppercase and lowercase alphanumeric characters, but the first character cannot be a number.

	Command	Purpose
Step 5	accept-lifetime <i>start-time</i> { infinite <i>end-time</i> duration <i>seconds</i> }	(Optional) Specify the time period during which the key can be received. The <i>start-time</i> and <i>end-time</i> syntax can be either <i>hh:mm:ss Month date year</i> or <i>hh:mm:ss date Month year</i> . The default is forever with the default <i>start-time</i> and the earliest acceptable date as January 1, 1993. The default <i>end-time</i> and duration is infinite .
Step 6	send-lifetime <i>start-time</i> { infinite <i>end-time</i> duration <i>seconds</i> }	(Optional) Specify the time period during which the key can be sent. The <i>start-time</i> and <i>end-time</i> syntax can be either <i>hh:mm:ss Month date year</i> or <i>hh:mm:ss date Month year</i> . The default is forever with the default <i>start-time</i> and the earliest acceptable date as January 1, 1993. The default <i>end-time</i> and duration is infinite .
Step 7	end	Return to privileged EXEC mode.
Step 8	show key chain	Display authentication key information.
Step 9	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To remove the key chain, use the **no key chain** *name-of-chain* global configuration command.

Monitoring and Maintaining the IP Network

You can remove all contents of a particular cache, table, or database. You can also display specific statistics. Use the privileged EXEC commands in [Table 32-16](#) to clear routes or display status:

Table 32-16 Commands to Clear IP Routes or Display Route Status

Command	Purpose
clear ip route { <i>network</i> [<i>mask</i> *] }	Clear one or more routes from the IP routing table.
show ip protocols	Display the parameters and state of the active routing protocol process.
show ip route [<i>address</i> [<i>mask</i>] [longer-prefixes]] [<i>protocol</i> [<i>process-id</i>]]	Display the current state of the routing table.
show ip route summary	Display the current state of the routing table in summary form.
show ip route supernets-only	Display supernets.
show ip cache	Display the routing table used to switch IP traffic.
show route-map [<i>map-name</i>]	Display all route maps configured or only the one specified.

