



Configuring IP Multicast Routing

This chapter describes how to configure IP multicast routing on the Catalyst 3750 switch. IP multicasting is a more efficient way to use network resources, especially for bandwidth-intensive services such as audio and video. IP multicast routing enables a host (source) to send packets to a group of hosts (receivers) anywhere within the IP network by using a special form of IP address called the IP *multicast group address*. The sending host inserts the multicast group address into the IP destination address field of the packet, and IP multicast routers and multilayer switches forward incoming IP multicast packets out all interfaces that lead to members of the multicast group. Any host, regardless of whether it is a member of a group, can send to a group. However, only the members of a group receive the message.

To use this feature, the stack master must be running the IP services image (formerly known as the enhanced multilayer image [EMI]). Unless otherwise noted, the term *switch* refers to a standalone switch and to a switch stack.



Note

For complete syntax and usage information for the commands used in this chapter, see the *Cisco IOS IP Command Reference, Volume 3 of 3: Multicast, Release 12.2*.

This chapter consists of these sections:

- [Understanding Cisco's Implementation of IP Multicast Routing, page 40-2](#)
- [Multicast Routing and Switch Stacks, page 40-8](#)
- [Configuring IP Multicast Routing, page 40-8](#)
- [Configuring Advanced PIM Features, page 40-23](#)
- [Configuring Optional IGMP Features, page 40-26](#)
- [Configuring Optional Multicast Routing Features, page 40-32](#)
- [Configuring Basic DVMRP Interoperability Features, page 40-36](#)
- [Configuring Advanced DVMRP Interoperability Features, page 40-41](#)
- [Monitoring and Maintaining IP Multicast Routing, page 40-49](#)

For information on configuring the Multicast Source Discovery Protocol (MSDP), see [Chapter 41, "Configuring MSDP."](#)

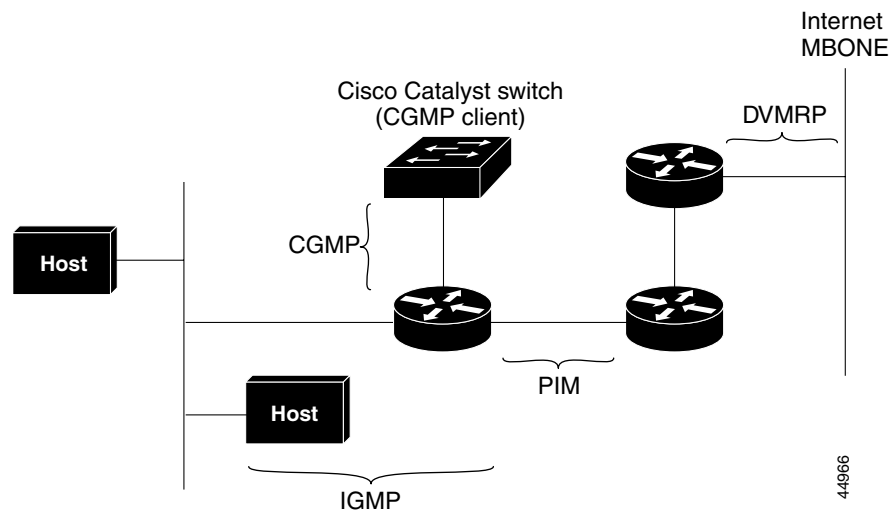
Understanding Cisco's Implementation of IP Multicast Routing

The Cisco IOS software supports these protocols to implement IP multicast routing:

- Internet Group Management Protocol (IGMP) is used among hosts on a LAN and the routers (and multilayer switches) on that LAN to track the multicast groups of which hosts are members.
- Protocol-Independent Multicast (PIM) protocol is used among routers and multilayer switches to track which multicast packets to forward to each other and to their directly connected LANs.
- Distance Vector Multicast Routing Protocol (DVMRP) is used on the multicast backbone of the Internet (MBONE). The software supports PIM-to-DVMRP interaction.
- Cisco Group Management Protocol (CGMP) is used on Cisco routers and multilayer switches connected to Layer 2 Catalyst switches to perform tasks similar to those performed by IGMP.

Figure 40-1 shows where these protocols operate within the IP multicast environment.

Figure 40-1 IP Multicast Routing Protocols



Understanding IGMP

To participate in IP multicasting, multicast hosts, routers, and multilayer switches must have the IGMP operating. This protocol defines the querier and host roles:

- A querier is a network device that sends query messages to discover which network devices are members of a given multicast group.
- A host is a receiver that sends report messages (in response to query messages) to inform a querier of a host membership.

A set of queriers and hosts that receive multicast data streams from the same source is called a multicast group. Queriers and hosts use IGMP messages to join and leave multicast groups.

Any host, regardless of whether it is a member of a group, can send to a group. However, only the members of a group receive the message. Membership in a multicast group is dynamic; hosts can join and leave at any time. There is no restriction on the location or number of members in a multicast group. A host can be a member of more than one multicast group at a time. How active a multicast group is and what members it has can vary from group to group and from time to time. A multicast group can be active for a long time, or it can be very short-lived. Membership in a group can constantly change. A group that has members can have no activity.

IP multicast traffic uses group addresses, which are class D addresses. The high-order bits of a Class D address are 1110. Therefore, host group addresses can be in the range 224.0.0.0 through 239.255.255.255. Multicast addresses in the range 224.0.0.0 to 224.0.0.255 are reserved for use by routing protocols and other network control traffic. The address 224.0.0.0 is guaranteed not to be assigned to any group.

IGMP packets are sent using these IP multicast group addresses:

- IGMP general queries are destined to the address 224.0.0.1 (all systems on a subnet).
- IGMP group-specific queries are destined to the group IP address for which the switch is querying.
- IGMP group membership reports are destined to the group IP address for which the switch is reporting.
- IGMP Version 2 (IGMPv2) leave messages are destined to the address 224.0.0.2 (all-multicast-routers on a subnet). In some old host IP stacks, leave messages might be destined to the group IP address rather than to the all-routers address.

IGMP Version 1

IGMP Version 1 (IGMPv1) primarily uses a query-response model that enables the multicast router and multilayer switch to find which multicast groups are active (have one or more hosts interested in a multicast group) on the local subnet. IGMPv1 has other processes that enable a host to join and leave a multicast group. For more information, see RFC 1112.

IGMP Version 2

IGMPv2 extends IGMP functionality by providing such features as the IGMP leave process to reduce leave latency, group-specific queries, and an explicit maximum query response time. IGMPv2 also adds the capability for routers to elect the IGMP querier without depending on the multicast protocol to perform this task. For more information, see RFC 2236.

Understanding PIM

PIM is called *protocol-independent*: regardless of the unicast routing protocols used to populate the unicast routing table, PIM uses this information to perform multicast forwarding instead of maintaining a separate multicast routing table.

PIM is defined in RFC 2362, *Protocol-Independent Multicast-Sparse Mode (PIM-SM): Protocol Specification*. PIM is defined in these Internet Engineering Task Force (IETF) Internet drafts:

- *Protocol Independent Multicast (PIM): Motivation and Architecture*
- *Protocol Independent Multicast (PIM), Dense Mode Protocol Specification*
- *Protocol Independent Multicast (PIM), Sparse Mode Protocol Specification*

- *draft-ietf-idmr-igmp-v2-06.txt, Internet Group Management Protocol, Version 2*
- *draft-ietf-pim-v2-dm-03.txt, PIM Version 2 Dense Mode*

PIM Versions

PIMv2 includes these improvements over PIMv1:

- A single, active rendezvous point (RP) exists per multicast group, with multiple backup RPs. This single RP compares to multiple active RPs for the same group in PIMv1.
- A bootstrap router (BSR) provides a fault-tolerant, automated RP discovery and distribution mechanism that enables routers and multilayer switches to dynamically learn the group-to-RP mappings.
- Sparse mode and dense mode are properties of a group, as opposed to an interface. We strongly recommend sparse-dense mode, as opposed to either sparse mode or dense mode only.
- PIM join and prune messages have more flexible encoding for multiple address families.
- A more flexible hello packet format replaces the query packet to encode current and future capability options.
- Register messages to an RP specify whether they are sent by a border router or a designated router.
- PIM packets are no longer inside IGMP packets; they are standalone packets.

PIM Modes

PIM can operate in dense mode (DM), sparse mode (SM), or in sparse-dense mode (PIM DM-SM), which handles both sparse groups and dense groups at the same time.

PIM DM

PIM DM builds source-based multicast distribution trees. In dense mode, a PIM DM router or multilayer switch assumes that all other routers or multilayer switches forward multicast packets for a group. If a PIM DM device receives a multicast packet and has no directly connected members or PIM neighbors present, a prune message is sent back to the source to stop unwanted multicast traffic. Subsequent multicast packets are not flooded to this router or switch on this pruned branch because branches without receivers are pruned from the distribution tree, leaving only branches that contain receivers.

When a new receiver on a previously pruned branch of the tree joins a multicast group, the PIM DM device detects the new receiver and immediately sends a graft message up the distribution tree toward the source. When the upstream PIM DM device receives the graft message, it immediately puts the interface on which the graft was received into the forwarding state so that the multicast traffic begins flowing to the receiver.

PIM SM

PIM SM uses shared trees and shortest-path-trees (SPTs) to distribute multicast traffic to multicast receivers in the network. In PIM SM, a router or multilayer switch assumes that other routers or switches do not forward multicast packets for a group, unless there is an explicit request for the traffic (join message). When a host joins a multicast group using IGMP, its directly connected PIM SM device sends PIM join messages toward the root, also known as the RP. This join message travels router-by-router toward the root, constructing a branch of the shared tree as it goes.

The RP keeps track of multicast receivers. It also registers sources through register messages received from the source's first-hop router (*designated router* [DR]) to complete the shared tree path from the source to the receiver. When using a shared tree, sources must send their traffic to the RP so that the traffic reaches all receivers.

Prune messages are sent up the distribution tree to prune multicast group traffic. This action permits branches of the shared tree or SPT that were created with explicit join messages to be torn down when they are no longer needed.

Auto-RP

This proprietary feature eliminates the need to manually configure the RP information in every router and multilayer switch in the network. For Auto-RP to work, you configure a Cisco router or multilayer switch as the mapping agent. It uses IP multicast to learn which routers or switches in the network are possible candidate RPs to receive candidate RP announcements. Candidate RPs periodically send multicast RP-announce messages to a particular group or group range to announce their availability.

Mapping agents listen to these candidate RP announcements and use the information to create entries in their Group-to-RP mapping caches. Only one mapping cache entry is created for any Group-to-RP range received, even if multiple candidate RPs are sending RP announcements for the same range. As the RP-announce messages arrive, the mapping agent selects the router or switch with the highest IP address as the active RP and stores this RP address in the Group-to-RP mapping cache.

Mapping agents periodically multicast the contents of their Group-to-RP mapping cache. Thus, all routers and switches automatically discover which RP to use for the groups they support. If a router or switch fails to receive RP-discovery messages and the Group-to-RP mapping information expires, it switches to a statically configured RP that was defined with the **ip pim rp-address** global configuration command. If no statically configured RP exists, the router or switch changes the group to dense-mode operation.

Multiple RPs serve different group ranges or serve as hot backups of each other.

Bootstrap Router

PIMv2 BSR is another method to distribute group-to-RP mapping information to all PIM routers and multilayer switches in the network. It eliminates the need to manually configure RP information in every router and switch in the network. However, instead of using IP multicast to distribute group-to-RP mapping information, BSR uses hop-by-hop flooding of special BSR messages to distribute the mapping information.

The BSR is elected from a set of candidate routers and switches in the domain that have been configured to function as BSRs. The election mechanism is similar to the root-bridge election mechanism used in bridged LANs. The BSR election is based on the BSR priority of the device contained in the BSR messages that are sent hop-by-hop through the network. Each BSR device examines the message and forwards out all interfaces only the message that has either a higher BSR priority than its BSR priority or the same BSR priority, but with a higher BSR IP address. Using this method, the BSR is elected.

The elected BSR sends BSR messages with a TTL of 1. Neighboring PIMv2 routers or multilayer switches receive the BSR message and multicast it out all other interfaces (except the one on which it was received) with a TTL of 1. In this way, BSR messages travel hop-by-hop throughout the PIM domain. Because BSR messages contain the IP address of the current BSR, the flooding mechanism enables candidate RPs to automatically learn which device is the elected BSR.

Candidate RPs send candidate RP advertisements showing the group range for which they are responsible to the BSR, which stores this information in its local candidate-RP cache. The BSR periodically advertises the contents of this cache in BSR messages to all other PIM devices in the domain. These messages travel hop-by-hop through the network to all routers and switches, which store the RP information in the BSR message in their local RP cache. The routers and switches select the same RP for a given group because they all use a common RP hashing algorithm.

Multicast Forwarding and Reverse Path Check

With unicast routing, routers and multilayer switches forward traffic through the network along a single path from the source to the destination host whose IP address appears in the destination address field of the IP packet. Each router and switch along the way makes a unicast forwarding decision, using the destination IP address in the packet, by looking up the destination address in the unicast routing table and forwarding the packet through the specified interface to the next hop toward the destination.

With multicasting, the source is sending traffic to an arbitrary group of hosts represented by a multicast group address in the destination address field of the IP packet. To decide whether to forward or drop an incoming multicast packet, the router or multilayer switch uses a reverse path forwarding (RPF) check on the packet as follows and shown in [Figure 40-2](#):

1. The router or multilayer switch examines the source address of the arriving multicast packet to decide whether the packet arrived on an interface that is on the reverse path back to the source.
2. If the packet arrives on the interface leading back to the source, the RPF check is successful and the packet is forwarded to all interfaces in the outgoing interface list (which might not be all interfaces on the router).
3. If the RPF check fails, the packet is discarded.

Some multicast routing protocols, such as DVMRP, maintain a separate multicast routing table and use it for the RPF check. However, PIM uses the unicast routing table to perform the RPF check.

[Figure 40-2](#) shows port 2 receiving a multicast packet from source 151.10.3.21. [Table 40-1](#) shows that the port on the reverse path to the source is port 1, not port 2. Because the RPF check fails, the multilayer switch discards the packet. Another multicast packet from source 151.10.3.21 is received on port 1, and the routing table shows this port is on the reverse path to the source. Because the RPF check passes, the switch forwards the packet to all port in the outgoing port list.

Figure 40-2 RPF Check

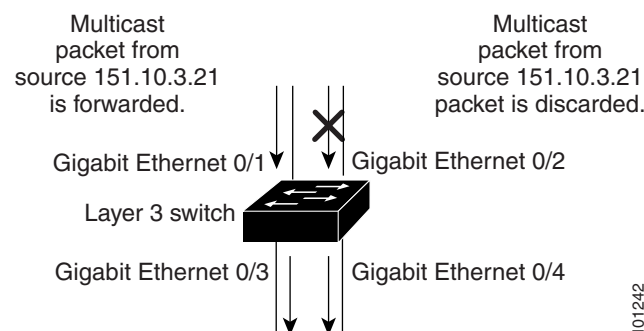


Table 40-1 **Routing Table Example for an RPF Check**

Network	Port
151.10.0.0/16	Gigabit Ethernet 1/0/1
198.14.32.0/32	Gigabit Ethernet 1/0/3
204.1.16.0/24	Gigabit Ethernet 1/0/4

PIM uses both source trees and RP-rooted shared trees to forward datagrams (described in the “[PIM DM](#)” section on page 40-4 and the “[PIM SM](#)” section on page 40-4). The RPF check is performed differently for each:

- If a PIM router or multilayer switch has a source-tree state (that is, an (S,G) entry is present in the multicast routing table), it performs the RPF check against the IP address of the source of the multicast packet.
- If a PIM router or multilayer switch has a shared-tree state (and no explicit source-tree state), it performs the RPF check on the RP address (which is known when members join the group).

Sparse-mode PIM uses the RPF lookup function to decide where it needs to send joins and prunes:

- (S,G) joins (which are source-tree states) are sent toward the source.
- (*,G) joins (which are shared-tree states) are sent toward the RP.

DVMRP and dense-mode PIM use only source trees and use RPF as previously described.

Understanding DVMRP

DVMRP is implemented in the equipment of many vendors and is based on the public-domain mroute program. This protocol has been deployed in the MBONE and in other intradomain multicast networks.

Cisco routers and multilayer switches run PIM and can forward multicast packets to and receive from a DVMRP neighbor. It is also possible to propagate DVMRP routes into and through a PIM cloud. The software propagates DVMRP routes and builds a separate database for these routes on each router and multilayer switch, but PIM uses this routing information to make the packet-forwarding decision. The software does not implement the complete DVMRP. However, it supports dynamic discovery of DVMRP routers and can interoperate with them over traditional media (such as Ethernet and FDDI) or over DVMRP-specific tunnels.

DVMRP neighbors build a route table by periodically exchanging source network routing information in route-report messages. The routing information stored in the DVMRP routing table is separate from the unicast routing table and is used to build a source distribution tree and to perform multicast forward using RPF.

DVMRP is a dense-mode protocol and builds a parent-child database using a constrained multicast model to build a forwarding tree rooted at the source of the multicast packets. Multicast packets are initially flooded down this source tree. If redundant paths are on the source tree, packets are not forwarded along those paths. Forwarding occurs until prune messages are received on those parent-child links, which further constrain the broadcast of multicast packets.

Understanding CGMP

This software release provides CGMP-server support on your switch; no client-side functionality is provided. The switch serves as a CGMP server for devices that do not support IGMP snooping but have CGMP-client functionality.

CGMP is a protocol used on Cisco routers and multilayer switches connected to Layer 2 Catalyst switches to perform tasks similar to those performed by IGMP. CGMP permits Layer 2 group membership information to be communicated from the CGMP server to the switch. The switch can then learn on which interfaces multicast members reside instead of flooding multicast traffic to all switch interfaces. (IGMP snooping is another method to constrain the flooding of multicast packets. For more information, see [Chapter 24, “Configuring IGMP Snooping and MVR.”](#))

CGMP is necessary because the Layer 2 switch cannot distinguish between IP multicast data packets and IGMP report messages, which are both at the MAC-level and are addressed to the same group address.

Multicast Routing and Switch Stacks

For all multicast routing protocols, the entire stack appears as a single router to the network and operates as a single multicast router.

In a Catalyst 3750 switch stack, the routing master (stack master) performs these functions:

- It is responsible for completing the IP multicast routing functions of the stack. It fully initializes and runs the IP multicast routing protocols.
- It builds and maintains the multicast routing table for the entire stack.
- It is responsible for distributing the multicast routing table to all stack members.

The stack members perform these functions:

- They act as multicast routing standby devices and are ready to take over if there is a stack master failure.

If the stack master fails, all stack members delete their multicast routing tables. The newly elected stack master starts building the routing tables and distributes them to the stack members.

**Note**

If a stack master running the IP services image fails and if the newly elected stack master is running the IP base image (formerly known as the standard multilayer image [SMI]), the switch stack will lose its multicast routing capability.

For information about the stack master election process, see [Chapter 5, “Managing Switch Stacks.”](#)

- They do not build multicast routing tables. Instead, they use the multicast routing table that is distributed by the stack master.

Configuring IP Multicast Routing

These sections contain this configuration information:

- [Default Multicast Routing Configuration, page 40-9](#)
- [Multicast Routing Configuration Guidelines, page 40-9](#)
- [Configuring Basic Multicast Routing, page 40-11](#) (required)

- [Configuring a Rendezvous Point, page 40-12](#) (required if the interface is in sparse-dense mode, and you want to treat the group as a sparse group)
- [Using Auto-RP and a BSR, page 40-22](#) (required for non-Cisco PIMv2 devices to interoperate with Cisco PIM v1 devices))
- [Monitoring the RP Mapping Information, page 40-23](#) (optional)
- [Troubleshooting PIMv1 and PIMv2 Interoperability Problems, page 40-23](#) (optional)

Default Multicast Routing Configuration

Table 40-2 shows the default multicast routing configuration.

Table 40-2 **Default Multicast Routing Configuration**

Feature	Default Setting
Multicast routing	Disabled on all interfaces.
PIM version	Version 2.
PIM mode	No mode is defined.
PIM RP address	None configured.
PIM domain border	Disabled.
PIM multicast boundary	None.
Candidate BSRs	Disabled.
Candidate RPs	Disabled.
Shortest-path tree threshold rate	0 kbps.
PIM router query message interval	30 seconds.

Multicast Routing Configuration Guidelines

To avoid misconfiguring multicast routing on your switch, review the information in these sections:

- [PIMv1 and PIMv2 Interoperability, page 40-9](#)
- [Auto-RP and BSR Configuration Guidelines, page 40-10](#)

PIMv1 and PIMv2 Interoperability

The Cisco PIMv2 implementation provides interoperability and transition between Version 1 and Version 2, although there might be some minor problems.

You can upgrade to PIMv2 incrementally. PIM Versions 1 and 2 can be configured on different routers and multilayer switches within one network. Internally, all routers and multilayer switches on a shared media network must run the same PIM version. Therefore, if a PIMv2 device detects a PIMv1 device, the Version 2 device downgrades itself to Version 1 until all Version 1 devices have been shut down or upgraded.

PIMv2 uses the BSR to discover and announce RP-set information for each group prefix to all the routers and multilayer switches in a PIM domain. PIMv1, together with the Auto-RP feature, can perform the same tasks as the PIMv2 BSR. However, Auto-RP is a standalone protocol, separate from PIMv1, and is a proprietary Cisco protocol. PIMv2 is a standards track protocol in the IETF. We recommend that you use PIMv2. The BSR mechanism interoperates with Auto-RP on Cisco routers and multilayer switches. For more information, see the [“Auto-RP and BSR Configuration Guidelines” section on page 40-10](#).

When PIMv2 devices interoperate with PIMv1 devices, Auto-RP should have already been deployed. A PIMv2 BSR that is also an Auto-RP mapping agent automatically advertises the RP elected by Auto-RP. That is, Auto-RP sets its single RP on every router or multilayer switch in the group. Not all routers and switches in the domain use the PIMv2 hash function to select multiple RPs.

Dense-mode groups in a mixed PIMv1 and PIMv2 region need no special configuration; they automatically interoperate.

Sparse-mode groups in a mixed PIMv1 and PIMv2 region are possible because the Auto-RP feature in PIMv1 interoperates with the PIMv2 RP feature. Although all PIMv2 devices can also use PIMv1, we recommend that the RPs be upgraded to PIMv2. To ease the transition to PIMv2, we have these recommendations:

- Use Auto-RP throughout the region.
- Configure sparse-dense mode throughout the region.

If Auto-RP is not already configured in the PIMv1 regions, configure Auto-RP. For more information, see the [“Configuring Auto-RP” section on page 40-14](#).

Auto-RP and BSR Configuration Guidelines

There are two approaches to using PIMv2. You can use Version 2 exclusively in your network or migrate to Version 2 by employing a mixed PIM version environment.

- If your network is all Cisco routers and multilayer switches, you can use either Auto-RP or BSR.
- If you have non-Cisco routers in your network, you must use BSR.
- If you have Cisco PIMv1 and PIMv2 routers and multilayer switches and non-Cisco routers, you must use both Auto-RP and BSR. If your network includes routers from other vendors, configure the Auto-RP mapping agent and the BSR on a Cisco PIMv2 device. Ensure that no PIMv1 device is located in the path a between the BSR and a non-Cisco PIMv2 device.
- Because bootstrap messages are sent hop-by-hop, a PIMv1 device prevents these messages from reaching all routers and multilayer switches in your network. Therefore, if your network has a PIMv1 device in it and only Cisco routers and multilayer switches, it is best to use Auto-RP.
- If you have a network that includes non-Cisco routers, configure the Auto-RP mapping agent and the BSR on a Cisco PIMv2 router or multilayer switch. Ensure that no PIMv1 device is on the path between the BSR and a non-Cisco PIMv2 router.
- If you have non-Cisco PIMv2 routers that need to interoperate with Cisco PIMv1 routers and multilayer switches, both Auto-RP and a BSR are required. We recommend that a Cisco PIMv2 device be both the Auto-RP mapping agent and the BSR. For more information, see the [“Using Auto-RP and a BSR” section on page 40-22](#).

Configuring Basic Multicast Routing

You must enable IP multicast routing and configure the PIM version and the PIM mode. Then the software can forward multicast packets, and the switch can populate its multicast routing table.

You can configure an interface to be in PIM dense mode, sparse mode, or sparse-dense mode. The switch populates its multicast routing table and forwards multicast packets it receives from its directly connected LANs according to the mode setting. You must enable PIM in one of these modes for an interface to perform IP multicast routing. Enabling PIM on an interface also enables IGMP operation on that interface.

In populating the multicast routing table, dense-mode interfaces are always added to the table. Sparse-mode interfaces are added to the table only when periodic join messages are received from downstream devices or when there is a directly connected member on the interface. When forwarding from a LAN, sparse-mode operation occurs if there is an RP known for the group. If so, the packets are encapsulated and sent toward the RP. When no RP is known, the packet is flooded in a dense-mode fashion. If the multicast traffic from a specific source is sufficient, the receiver's first-hop router might send join messages toward the source to build a source-based distribution tree.

By default, multicast routing is disabled, and there is no default mode setting. This procedure is required.

Beginning in privileged EXEC mode, follow these steps to enable IP multicasting, to configure a PIM version, and to configure a PIM mode. This procedure is required.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip multicast-routing distributed	Enable IP multicast distributed switching.
Step 3	interface <i>interface-id</i>	Specify the Layer 3 interface on which you want to enable multicast routing, and enter interface configuration mode. The specified interface must be one of the following: - A routed port: a physical port that has been configured as a Layer 3 port by entering the no switchport interface configuration command. - An SVI: a VLAN interface created by using the interface vlan <i>vlan-id</i> global configuration command. These interfaces must have IP addresses assigned to them. For more information, see the “Configuring Layer 3 Interfaces” section on page 11-25.
Step 4	ip pim version [1 2]	Configure the PIM version on the interface. By default, Version 2 is enabled and is the recommended setting. An interface in PIMv2 mode automatically downgrades to PIMv1 mode if that interface has a PIMv1 neighbor. The interface returns to Version 2 mode after all Version 1 neighbors are shut down or upgraded. For more information, see the “PIMv1 and PIMv2 Interoperability” section on page 40-9.

	Command	Purpose
Step 5	ip pim {dense-mode sparse-mode sparse-dense-mode}	Enable a PIM mode on the interface. By default, no mode is configured. The keywords have these meanings: • dense-mode—Enables dense mode of operation. • sparse-mode—Enables sparse mode of operation. If you configure sparse-mode, you must also configure an RP. For more information, see the “Configuring a Rendezvous Point” section on page 40-12. • sparse-dense-mode—Causes the interface to be treated in the mode in which the group belongs. Sparse-dense-mode is the recommended setting. Note After you enable a PIM mode on the interface, the ip mroute-cache distributed interface configuration command is automatically entered for the interface and appears in the running configuration.
Step 6	end	Return to privileged EXEC mode.
Step 7	show running-config	Verify your entries.
Step 8	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To disable multicasting, use the **no ip multicast-routing distributed** global configuration command. To return to the default PIM version, use the **no ip pim version** interface configuration command. To disable PIM on an interface, use the **no ip pim** interface configuration command.

Configuring a Rendezvous Point

You must have an RP if the interface is in sparse-dense mode and if you want to treat the group as a sparse group. You can use several methods, as described in these sections:

- [Manually Assigning an RP to Multicast Groups](#), page 40-12
- [Configuring Auto-RP](#), page 40-14 (a standalone, Cisco-proprietary protocol separate from PIMv1)
- [Configuring PIMv2 BSR](#), page 40-18 (a standards track protocol in the Internet Engineering Task Force (IETF))

You can use Auto-RP, BSR, or a combination of both, depending on the PIM version you are running and the types of routers in your network. For more information, see the [“PIMv1 and PIMv2 Interoperability”](#) section on page 40-9 and the [“Auto-RP and BSR Configuration Guidelines”](#) section on page 40-10.

Manually Assigning an RP to Multicast Groups

This section explains how to manually configure an RP. If the RP for a group is learned through a dynamic mechanism (such as Auto-RP or BSR), you need not perform this task for that RP.

Senders of multicast traffic announce their existence through register messages received from the source’s first-hop router (designated router) and forwarded to the RP. Receivers of multicast packets use RPs to join a multicast group by using explicit join messages. RPs are not members of the multicast group; rather, they serve as a *meeting place* for multicast sources and group members.

You can configure a single RP for multiple groups defined by an access list. If there is no RP configured for a group, the multilayer switch treats the group as dense and uses the dense-mode PIM techniques.

Beginning in privileged EXEC mode, follow these steps to manually configure the address of the RP. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip pim rp-address <i>ip-address</i> [<i>access-list-number</i>] [override]	Configure the address of a PIM RP. By default, no PIM RP address is configured. You must configure the IP address of RPs on all routers and multilayer switches (including the RP). If there is no RP configured for a group, the switch treats the group as dense, using the dense-mode PIM techniques. A PIM device can be an RP for more than one group. Only one RP address can be used at a time within a PIM domain. The access-list conditions specify for which groups the device is an RP. - For <i>ip-address</i>, enter the unicast address of the RP in dotted-decimal notation. (Optional) For <i>access-list-number</i>, enter an IP standard access list number from 1 to 99. If no access list is configured, the RP is used for all groups. (Optional) The override keyword means that if there is a conflict between the RP configured with this command and one learned by Auto-RP or BSR, the RP configured with this command prevails.
Step 3	access-list <i>access-list-number</i> { deny permit } <i>source</i> [<i>source-wildcard</i>]	Create a standard access list, repeating the command as many times as necessary. - For <i>access-list-number</i>, enter the access list number specified in Step 2. - The deny keyword denies access if the conditions are matched. The permit keyword permits access if the conditions are matched. - For <i>source</i>, enter the multicast group address for which the RP should be used. (Optional) For <i>source-wildcard</i>, enter the wildcard bits in dotted decimal notation to be applied to the source. Place ones in the bit positions that you want to ignore. Recall that the access list is always terminated by an implicit deny statement for everything.
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 remove an RP address, use the **no ip pim rp-address** *ip-address* [*access-list-number*] [**override**] global configuration command.

This example shows how to configure the address of the RP to 147.106.6.22 for multicast group 225.2.2.2 only:

```
Switch(config)# access-list 1 permit 225.2.2.2 0.0.0.0
Switch(config)# ip pim rp-address 147.106.6.22 1
```

Configuring Auto-RP

Auto-RP uses IP multicast to automate the distribution of group-to-RP mappings to all Cisco routers and multilayer switches in a PIM network. It has these benefits:

- It is easy to use multiple RPs within a network to serve different group ranges.
- It provides load splitting among different RPs and arrangement of RPs according to the location of group participants.
- It avoids inconsistent, manual RP configurations on every router and multilayer switch in a PIM network, which can cause connectivity problems.



Note

If you configure PIM in sparse mode or sparse-dense mode and do not configure Auto-RP, you must manually configure an RP as described in the [“Manually Assigning an RP to Multicast Groups” section on page 40-12](#).



Note

If routed interfaces are configured in sparse mode, Auto-RP can still be used if all devices are configured with a manual RP address for the Auto-RP groups.

These sections describe how to configure Auto-RP:

- [Setting up Auto-RP in a New Internetwork, page 40-14](#) (optional)
- [Adding Auto-RP to an Existing Sparse-Mode Cloud, page 40-15](#) (optional)
- [Preventing Join Messages to False RPs, page 40-16](#) (optional)
- [Filtering Incoming RP Announcement Messages, page 40-17](#) (optional)

For overview information, see the [“Auto-RP” section on page 40-5](#).

Setting up Auto-RP in a New Internetwork

If you are setting up Auto-RP in a new internetwork, you do not need a default RP because you configure all the interfaces for sparse-dense mode. Follow the process described in the [“Adding Auto-RP to an Existing Sparse-Mode Cloud” section on page 40-15](#). However, omit Step 3 if you want to configure a PIM router as the RP for the local group.

Adding Auto-RP to an Existing Sparse-Mode Cloud

This section contains some suggestions for the initial deployment of Auto-RP into an existing sparse-mode cloud to minimize disruption of the existing multicast infrastructure.

Beginning in privileged EXEC mode, follow these steps to deploy Auto-RP in an existing sparse-mode cloud. This procedure is optional.

	Command	Purpose
Step 1	show running-config	Verify that a default RP is already configured on all PIM devices and the RP in the sparse-mode network. It was previously configured with the ip pim rp-address global configuration command. This step is not required for sparse-dense-mode environments. The selected RP should have good connectivity and be available across the network. Use this RP for the global groups (for example 224.x.x.x and other global groups). Do not reconfigure the group address range that this RP serves. RPs dynamically discovered through Auto-RP take precedence over statically configured RPs. Assume that it is desirable to use a second RP for the local groups.
Step 2	configure terminal	Enter global configuration mode.
Step 3	ip pim send-rp-announce <i>interface-id</i> scope <i>ttl</i> group-list <i>access-list-number</i> interval <i>seconds</i>	Configure another PIM device to be the candidate RP for local groups. - For <i>interface-id</i>, enter the interface type and number that identifies the RP address. Valid interfaces include physical ports, port channels, and VLANs. - For scope <i>ttl</i>, specify the time-to-live value in hops. Enter a hop count that is high enough so that the RP-announce messages reach all mapping agents in the network. There is no default setting. The range is 1 to 255. - For group-list <i>access-list-number</i>, enter an IP standard access list number from 1 to 99. If no access list is configured, the RP is used for all groups. - For interval <i>seconds</i>, specify how often the announcement messages must be sent. The default is 60 seconds. The range is 1 to 16383.
Step 4	access-list <i>access-list-number</i> { deny permit } <i>source</i> [<i>source-wildcard</i>]	Create a standard access list, repeating the command as many times as necessary. - For <i>access-list-number</i>, enter the access list number specified in Step 3. - The deny keyword denies access if the conditions are matched. The permit keyword permits access if the conditions are matched. - For <i>source</i>, enter the multicast group address range for which the RP should be used. (Optional) For <i>source-wildcard</i>, enter the wildcard bits in dotted decimal notation to be applied to the source. Place ones in the bit positions that you want to ignore. Recall that the access list is always terminated by an implicit deny statement for everything.

	Command	Purpose
Step 5	ip pim send-rp-discovery scope ttl	Find a switch whose connectivity is not likely to be interrupted, and assign it the role of RP-mapping agent. For scope ttl , specify the time-to-live value in hops to limit the RP discovery packets. All devices within the hop count from the source device receive the Auto-RP discovery messages. These messages tell other devices which group-to-RP mapping to use to avoid conflicts (such as overlapping group-to-RP ranges). There is no default setting. The range is 1 to 255.
Step 6	end	Return to privileged EXEC mode.
Step 7	show running-config	Verify your entries.
	show ip pim rp mapping	Display active RPs that are cached with associated multicast routing entries.
	show ip pim rp	Display the information cached in the routing table.
Step 8	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To remove the PIM device configured as the candidate RP, use the **no ip pim send-rp-announce interface-id** global configuration command. To remove the switch as the RP-mapping agent, use the **no ip pim send-rp-discovery** global configuration command.

This example shows how to send RP announcements out all PIM-enabled interfaces for a maximum of 31 hops. The IP address of port 1 is the RP. Access list 5 describes the group for which this switch serves as RP:

```
Switch(config)# ip pim send-rp-announce gigabitethernet1/0/1 scope 31 group-list 5
Switch(config)# access-list 5 permit 224.0.0.0 15.255.255.255
```

Preventing Join Messages to False RPs

Find whether the **ip pim accept-rp** command was previously configured throughout the network by using the **show running-config** privileged EXEC command. If the **ip pim accept-rp** command is not configured on any device, this problem can be addressed later. In those routers or multilayer switches already configured with the **ip pim accept-rp** command, you must enter the command again to accept the newly advertised RP.

To accept all RPs advertised with Auto-RP and reject all other RPs by default, use the **ip pim accept-rp auto-rp** global configuration command. This procedure is optional.

If all interfaces are in sparse mode, use a default-configured RP to support the two well-known groups 224.0.1.39 and 224.0.1.40. Auto-RP uses these two well-known groups to collect and distribute RP-mapping information. When this is the case and the **ip pim accept-rp auto-rp** command is configured, another **ip pim accept-rp** command accepting the RP must be configured as follows:

```
Switch(config)# ip pim accept-rp 172.10.20.1 1
Switch(config)# access-list 1 permit 224.0.1.39
Switch(config)# access-list 1 permit 224.0.1.40
```

Filtering Incoming RP Announcement Messages

You can add configuration commands to the mapping agents to prevent a maliciously configured router from masquerading as a candidate RP and causing problems.

Beginning in privileged EXEC mode, follow these steps to filter incoming RP announcement messages. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip pim rp-announce-filter rp-list access-list-number group-list access-list-number	Filter incoming RP announcement messages. Enter this command on each mapping agent in the network. Without this command, all incoming RP-announce messages are accepted by default. For rp-list access-list-number, configure an access list of candidate RP addresses that, if permitted, is accepted for the group ranges supplied in the group-list access-list-number variable. If this variable is omitted, the filter applies to all multicast groups. If more than one mapping agent is used, the filters must be consistent across all mapping agents to ensure that no conflicts occur in the Group-to-RP mapping information.
Step 3	access-list access-list-number {deny permit} source [source-wildcard]	Create a standard access list, repeating the command as many times as necessary. - For <i>access-list-number</i>, enter the access list number specified in Step 2. - The deny keyword denies access if the conditions are matched. The permit keyword permits access if the conditions are matched. - Create an access list that specifies from which routers and multilayer switches the mapping agent accepts candidate RP announcements (rp-list ACL). - Create an access list that specifies the range of multicast groups from which to accept or deny (group-list ACL). - For <i>source</i>, enter the multicast group address range for which the RP should be used. (Optional) For <i>source-wildcard</i>, enter the wildcard bits in dotted decimal notation to be applied to the source. Place ones in the bit positions that you want to ignore. Recall that the access list is always terminated by an implicit deny statement for everything.
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 remove a filter on incoming RP announcement messages, use the **no ip pim rp-announce-filter rp-list access-list-number [group-list access-list-number]** global configuration command.

This example shows a sample configuration on an Auto-RP mapping agent that is used to prevent candidate RP announcements from being accepted from unauthorized candidate RPs:

```
Switch(config)# ip pim rp-announce-filter rp-list 10 group-list 20
Switch(config)# access-list 10 permit host 172.16.5.1
Switch(config)# access-list 10 permit host 172.16.2.1
Switch(config)# access-list 20 deny 239.0.0.0 0.0.255.255
Switch(config)# access-list 20 permit 224.0.0.0 15.255.255.255
```

In this example, the mapping agent accepts candidate RP announcements from only two devices, 172.16.5.1 and 172.16.2.1. The mapping agent accepts candidate RP announcements from these two devices only for multicast groups that fall in the group range of 224.0.0.0 to 239.255.255.255. The mapping agent does not accept candidate RP announcements from any other devices in the network. Furthermore, the mapping agent does not accept candidate RP announcements from 172.16.5.1 or 172.16.2.1 if the announcements are for any groups in the 239.0.0.0 through 239.255.255.255 range. This range is the administratively scoped address range.

Configuring PIMv2 BSR

These sections describe how to set up BSR in your PIMv2 network:

- [Defining the PIM Domain Border, page 40-18](#) (optional)
- [Defining the IP Multicast Boundary, page 40-19](#) (optional)
- [Configuring Candidate BSRs, page 40-20](#) (optional)
- [Configuring Candidate RPs, page 40-21](#) (optional)

For overview information, see the “Bootstrap Router” section on page 40-5.

Defining the PIM Domain Border

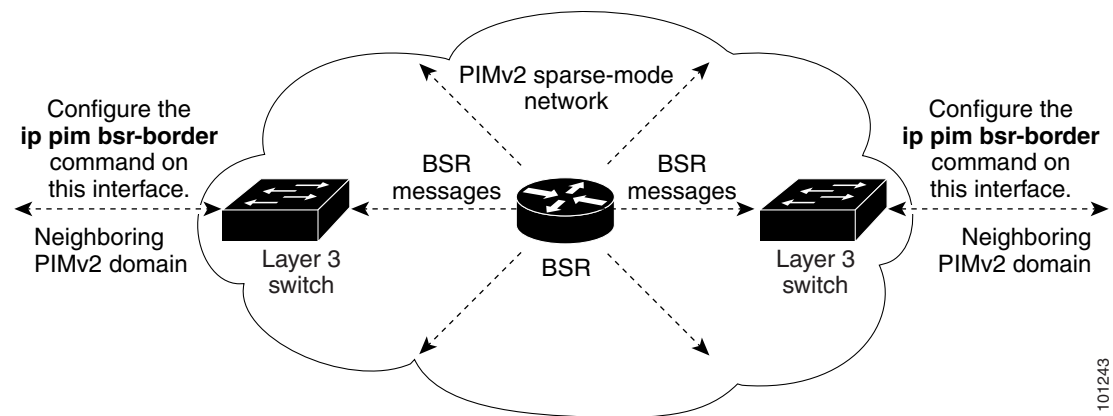
As IP multicast becomes more widespread, the chance of one PIMv2 domain bordering another PIMv2 domain is increasing. Because these two domains probably do not share the same set of RPs, BSR, candidate RPs, and candidate BSRs, you need to constrain PIMv2 BSR messages from flowing into or out of the domain. Allowing these messages to leak across the domain borders could adversely affect the normal BSR election mechanism and elect a single BSR across all bordering domains and co-mingle candidate RP advertisements, resulting in the election of RPs in the wrong domain.

Beginning in privileged EXEC mode, follow these steps to define the PIM domain border. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 3	ip pim bsr-border	Define a PIM bootstrap message boundary for the PIM domain. Enter this command on each interface that connects to other bordering PIM domains. This command instructs the switch to neither send or receive PIMv2 BSR messages on this interface as shown in Figure 40-3 .
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 remove the PIM border, use the **no ip pim bsr-border** interface configuration command.

Figure 40-3 Constraining PIMv2 BSR Messages



101243

Defining the IP Multicast Boundary

You define a multicast boundary to prevent Auto-RP messages from entering the PIM domain. You create an access list to deny packets destined for 224.0.1.39 and 224.0.1.40, which carry Auto-RP information.

Beginning in privileged EXEC mode, follow these steps to define a multicast boundary. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	access-list <i>access-list-number</i> deny <i>source</i> [<i>source-wildcard</i>]	Create a standard access list, repeating the command as many times as necessary. - For <i>access-list-number</i>, the range is 1 to 99. - The deny keyword denies access if the conditions are matched. - For <i>source</i>, enter multicast addresses 224.0.1.39 and 224.0.1.40, which carry Auto-RP information. (Optional) For <i>source-wildcard</i>, enter the wildcard bits in dotted decimal notation to be applied to the source. Place ones in the bit positions that you want to ignore. Recall that the access list is always terminated by an implicit deny statement for everything.
Step 3	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 4	ip multicast boundary <i>access-list-number</i>	Configure the boundary, specifying the access list you created in Step 2.
Step 5	end	Return to privileged EXEC mode.
Step 6	show running-config	Verify your entries.
Step 7	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To remove the boundary, use the **no ip multicast boundary** interface configuration command.

This example shows a portion of an IP multicast boundary configuration that denies Auto-RP information:

```
Switch(config)# access-list 1 deny 224.0.1.39
Switch(config)# access-list 1 deny 224.0.1.40
Switch(config)# interface gigabitethernet1/0/1
Switch(config-if)# ip multicast boundary 1
```

Configuring Candidate BSRs

You can configure one or more candidate BSRs. The devices serving as candidate BSRs should have good connectivity to other devices and be in the backbone portion of the network.

Beginning in privileged EXEC mode, follow these steps to configure your switch as a candidate BSR. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip pim bsr-candidate <i>interface-id</i> <i>hash-mask-length</i> [<i>priority</i>]	Configure your switch to be a candidate BSR. - For <i>interface-id</i>, enter the interface on this switch from which the BSR address is derived to make it a candidate. This interface must be enabled with PIM. Valid interfaces include physical ports, port channels, and VLANs. - For <i>hash-mask-length</i>, specify the mask length (32 bits maximum) that is to be ANDed with the group address before the hash function is called. All groups with the same seed hash correspond to the same RP. For example, if this value is 24, only the first 24 bits of the group addresses matter. (Optional) For <i>priority</i>, enter a number from 0 to 255. The BSR with the larger priority is preferred. If the priority values are the same, the device with the highest IP address is selected as the BSR. The default is 0.
Step 3	end	Return to privileged EXEC mode.
Step 4	show running-config	Verify your entries.
Step 5	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To remove this device as a candidate BSR, use the **no ip pim bsr-candidate** global configuration command.

This example shows how to configure a candidate BSR, which uses the IP address 172.21.24.18 on a port as the advertised BSR address, uses 30 bits as the hash-mask-length, and has a priority of 10.

```
Switch(config)# interface gigabitethernet1/0/2
Switch(config-if)# ip address 172.21.24.18 255.255.255.0
Switch(config-if)# ip pim sparse-dense-mode
Switch(config-if)# ip pim bsr-candidate gigabitethernet1/0/2 30 10
```

Configuring Candidate RPs

You can configure one or more candidate RPs. Similar to BSRs, the RPs should also have good connectivity to other devices and be in the backbone portion of the network. An RP can serve the entire IP multicast address space or a portion of it. Candidate RPs send candidate RP advertisements to the BSR. When deciding which devices should be RPs, consider these options:

- In a network of Cisco routers and multilayer switches where only Auto-RP is used, any device can be configured as an RP.
- In a network that includes only Cisco PIMv2 routers and multilayer switches and with routers from other vendors, any device can be used as an RP.
- In a network of Cisco PIMv1 routers, Cisco PIMv2 routers, and routers from other vendors, configure only Cisco PIMv2 routers and multilayer switches as RPs.

Beginning in privileged EXEC mode, follow these steps to configure your switch to advertise itself as a PIMv2 candidate RP to the BSR. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip pim rp-candidate <i>interface-id</i> [group-list <i>access-list-number</i>]	Configure your switch to be a candidate RP. • For <i>interface-id</i>, specify the interface whose associated IP address is advertised as a candidate RP address. Valid interfaces include physical ports, port channels, and VLANs. • (Optional) For group-list <i>access-list-number</i>, enter an IP standard access list number from 1 to 99. If no group-list is specified, the switch is a candidate RP for all groups.
Step 3	access-list <i>access-list-number</i> { deny permit } <i>source</i> [<i>source-wildcard</i>]	Create a standard access list, repeating the command as many times as necessary. • For <i>access-list-number</i>, enter the access list number specified in Step 2. • The deny keyword denies access if the conditions are matched. The permit keyword permits access if the conditions are matched. • For <i>source</i>, enter the number of the network or host from which the packet is being sent. • (Optional) For <i>source-wildcard</i>, enter the wildcard bits in dotted decimal notation to be applied to the source. Place ones in the bit positions that you want to ignore. Recall that the access list is always terminated by an implicit deny statement for everything.
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 remove this device as a candidate RP, use the **no ip pim rp-candidate** *interface-id* global configuration command.

This example shows how to configure the switch to advertise itself as a candidate RP to the BSR in its PIM domain. Standard access list number 4 specifies the group prefix associated with the RP that has the address identified by a port. That RP is responsible for the groups with the prefix 239.

```
Switch(config)# ip pim rp-candidate gigabitethernet1/0/2 group-list 4
Switch(config)# access-list 4 permit 239.0.0.0 0.255.255.255
```

Using Auto-RP and a BSR

If there are only Cisco devices in you network (no routers from other vendors), there is no need to configure a BSR. Configure Auto-RP in a network that is running both PIMv1 and PIMv2.

If you have non-Cisco PIMv2 routers that need to interoperate with Cisco PIMv1 routers and multilayer switches, both Auto-RP and a BSR are required. We recommend that a Cisco PIMv2 router or multilayer switch be both the Auto-RP mapping agent and the BSR.

If you must have one or more BSRs, we have these recommendations:

- Configure the candidate BSRs as the RP-mapping agents for Auto-RP. For more information, see the [“Configuring Auto-RP” section on page 40-14](#) and the [“Configuring Candidate BSRs” section on page 40-20](#).
- For group prefixes advertised through Auto-RP, the PIMv2 BSR mechanism should not advertise a subrange of these group prefixes served by a different set of RPs. In a mixed PIMv1 and PIMv2 domain, have backup RPs serve the same group prefixes. This prevents the PIMv2 DRs from selecting a different RP from those PIMv1 DRs, due to the longest match lookup in the RP-mapping database.

Beginning in privileged EXEC mode, follow these steps to verify the consistency of group-to-RP mappings. This procedure is optional.

	Command	Purpose
Step 1	<code>show ip pim rp [[group-name group-address] mapping]</code>	On any Cisco device, display the available RP mappings. • (Optional) For <i>group-name</i>, specify the name of the group about which to display RPs. • (Optional) For <i>group-address</i>, specify the address of the group about which to display RPs. • (Optional) Use the mapping keyword to display all group-to-RP mappings of which the Cisco device is aware (either configured or learned from Auto-RP).
Step 2	<code>show ip pim rp-hash group</code>	On a PIMv2 router or multilayer switch, confirm that the same RP is the one that a PIMv1 system chooses. For <i>group</i>, enter the group address for which to display RP information.

Monitoring the RP Mapping Information

To monitor the RP mapping information, use these commands in privileged EXEC mode:

- **show ip pim bsr** displays information about the elected BSR.
- **show ip pim rp-hash** *group* displays the RP that was selected for the specified group.
- **show ip pim rp** [*group-name* | *group-address* | **mapping**] displays how the switch learns of the RP (through the BSR or the Auto-RP mechanism).

Troubleshooting PIMv1 and PIMv2 Interoperability Problems

When debugging interoperability problems between PIMv1 and PIMv2, check these in the order shown:

1. Verify RP mapping with the **show ip pim rp-hash** privileged EXEC command, making sure that all systems agree on the same RP for the same group.
2. Verify interoperability between different versions of DRs and RPs. Make sure the RPs are interacting with the DRs properly (by responding with register-stops and forwarding decapsulated data packets from registers).

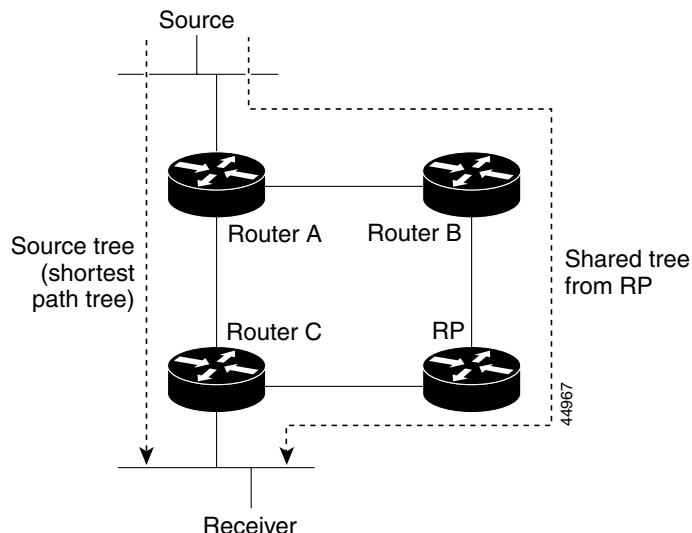
Configuring Advanced PIM Features

These sections describe the optional advanced PIM features:

- [Understanding PIM Shared Tree and Source Tree, page 40-23](#)
- [Delaying the Use of PIM Shortest-Path Tree, page 40-25](#) (optional)
- [Modifying the PIM Router-Query Message Interval, page 40-26](#) (optional)

Understanding PIM Shared Tree and Source Tree

By default, members of a group receive data from senders to the group across a single data-distribution tree rooted at the RP. [Figure 40-4](#) shows this type of shared-distribution tree. Data from senders is delivered to the RP for distribution to group members joined to the shared tree.

Figure 40-4 Shared Tree and Source Tree (Shortest-Path Tree)

If the data rate warrants, leaf routers (routers without any downstream connections) on the shared tree can use the data distribution tree rooted at the source. This type of distribution tree is called a shortest-path tree or source tree. By default, the software switches to a source tree upon receiving the first data packet from a source.

This process describes the move from a shared tree to a source tree:

1. A receiver joins a group; leaf Router C sends a join message toward the RP.
2. The RP puts a link to Router C in its outgoing interface list.
3. A source sends data; Router A encapsulates the data in a register message and sends it to the RP.
4. The RP forwards the data down the shared tree to Router C and sends a join message toward the source. At this point, data might arrive twice at Router C, once encapsulated and once natively.
5. When data arrives natively (unencapsulated) at the RP, it sends a register-stop message to Router A.
6. By default, reception of the first data packet prompts Router C to send a join message toward the source.
7. When Router C receives data on (S,G), it sends a prune message for the source up the shared tree.
8. The RP deletes the link to Router C from the outgoing interface of (S,G). The RP triggers a prune message toward the source.

Join and prune messages are sent for sources and RPs. They are sent hop-by-hop and are processed by each PIM device along the path to the source or RP. Register and register-stop messages are not sent hop-by-hop. They are sent by the designated router that is directly connected to a source and are received by the RP for the group.

Multiple sources sending to groups use the shared tree.

You can configure the PIM device to stay on the shared tree. For more information, see the [“Delaying the Use of PIM Shortest-Path Tree”](#) section on page 40-25.

Delaying the Use of PIM Shortest-Path Tree

The change from shared to source tree happens when the first data packet arrives at the last-hop router (Router C in [Figure 40-4](#)). This change occurs because the **ip pim spt-threshold** global configuration command controls that timing.

The shortest-path tree requires more memory than the shared tree but reduces delay. You might want to postpone its use. Instead of allowing the leaf router to immediately move to the shortest-path tree, you can specify that the traffic must first reach a threshold.

You can configure when a PIM leaf router should join the shortest-path tree for a specified group. If a source sends at a rate greater than or equal to the specified kbps rate, the multilayer switch triggers a PIM join message toward the source to construct a source tree (shortest-path tree). If the traffic rate from the source drops below the threshold value, the leaf router switches back to the shared tree and sends a prune message toward the source.

You can specify to which groups the shortest-path tree threshold applies by using a group list (a standard access list). If a value of 0 is specified or if the group list is not used, the threshold applies to all groups.

Beginning in privileged EXEC mode, follow these steps to configure a traffic rate threshold that must be reached before multicast routing is switched from the source tree to the shortest-path tree. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	access-list <i>access-list-number</i> { deny permit } <i>source</i> [<i>source-wildcard</i>]	Create a standard access list. - For <i>access-list-number</i>, the range is 1 to 99. - The deny keyword denies access if the conditions are matched. The permit keyword permits access if the conditions are matched. - For <i>source</i>, specify the multicast group to which the threshold will apply. (Optional) For <i>source-wildcard</i>, enter the wildcard bits in dotted decimal notation to be applied to the source. Place ones in the bit positions that you want to ignore. Recall that the access list is always terminated by an implicit deny statement for everything.
Step 3	ip pim spt-threshold { <i>kbps</i> infinity } [group-list <i>access-list-number</i>]	Specify the threshold that must be reached before moving to shortest-path tree (spt). - For <i>kbps</i>, specify the traffic rate in kilobits per second. The default is 0 kbps. Note Because of switch hardware limitations, 0 kbps is the only valid entry even though the range is 0 to 4294967. - Specify infinity if you want all sources for the specified group to use the shared tree, never switching to the source tree. (Optional) For group-list <i>access-list-number</i>, specify the access list created in Step 2. If the value is 0 or if the group-list is not used, the threshold applies to all groups.

	Command	Purpose
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 setting, use the **no ip pim spt-threshold** {*kpbs* | **infinity**} global configuration command.

Modifying the PIM Router-Query Message Interval

PIM routers and multilayer switches send PIM router-query messages to find which device will be the DR for each LAN segment (subnet). The DR is responsible for sending IGMP host-query messages to all hosts on the directly connected LAN.

With PIM DM operation, the DR has meaning only if IGMPv1 is in use. IGMPv1 does not have an IGMP querier election process, so the elected DR functions as the IGMP querier. With PIM SM operation, the DR is the device that is directly connected to the multicast source. It sends PIM register messages to notify the RP that multicast traffic from a source needs to be forwarded down the shared tree. In this case, the DR is the device with the highest IP address.

Beginning in privileged EXEC mode, follow these steps to modify the router-query message interval. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 3	ip pim query-interval <i>seconds</i>	Configure the frequency at which the switch sends PIM router-query messages. The default is 30 seconds. The range is 1 to 65535.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip igmp interface [<i>interface-id</i>]	Verify your entries.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To return to the default setting, use the **no ip pim query-interval** [*seconds*] interface configuration command.

Configuring Optional IGMP Features

These sections contain this configuration information:

- [Default IGMP Configuration, page 40-27](#)
- [Configuring the Switch as a Member of a Group, page 40-27](#) (optional)
- [Controlling Access to IP Multicast Groups, page 40-28](#) (optional)
- [Changing the IGMP Version, page 40-29](#) (optional)

- [Modifying the IGMP Host-Query Message Interval, page 40-29](#) (optional)
- [Changing the IGMP Query Timeout for IGMPv2, page 40-30](#) (optional)
- [Changing the Maximum Query Response Time for IGMPv2, page 40-31](#) (optional)
- [Configuring the Switch as a Statically Connected Member, page 40-31](#) (optional)

Default IGMP Configuration

Table 40-3 shows the default IGMP configuration.

Table 40-3 **Default IGMP Configuration**

Feature	Default Setting
Multilayer switch as a member of a multicast group	No group memberships are defined.
Access to multicast groups	All groups are allowed on an interface.
IGMP version	Version 2 on all interfaces.
IGMP host-query message interval	60 seconds on all interfaces.
IGMP query timeout	60 seconds on all interfaces.
IGMP maximum query response time	10 seconds on all interfaces.
Multilayer switch as a statically connected member	Disabled.

Configuring the Switch as a Member of a Group

You can configure the switch as a member of a multicast group and discover multicast reachability in a network. If all the multicast-capable routers and multilayer switches that you administer are members of a multicast group, pinging that group causes all these devices to respond. The devices respond to ICMP echo-request packets addressed to a group of which they are members. Another example is the multicast trace-route tools provided in the software.



Caution

Performing this procedure might impact the CPU performance because the CPU will receive all data traffic for the group address.

Beginning in privileged EXEC mode, follow these steps to configure the switch to be a member of a group. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 3	ip igmp join-group <i>group-address</i>	Configure the switch to join a multicast group. By default, no group memberships are defined. For <i>group-address</i> , specify the multicast IP address in dotted decimal notation.

	Command	Purpose
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip igmp interface <i>[interface-id]</i>	Verify your entries.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To cancel membership in a group, use the **no ip igmp join-group** *group-address* interface configuration command.

This example shows how to enable the switch to join multicast group 255.2.2.2:

```
Switch(config)# interface gigabitethernet1/0/1
Switch(config-if)# ip igmp join-group 255.2.2.2
```

Controlling Access to IP Multicast Groups

The switch sends IGMP host-query messages to find which multicast groups have members on attached local networks. The switch then forwards to these group members all packets addressed to the multicast group. You can place a filter on each interface to restrict the multicast groups that hosts on the subnet serviced by the interface can join.

Beginning in privileged EXEC mode, follow these steps to filter multicast groups allowed on an interface. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 3	ip igmp access-group <i>access-list-number</i>	Specify the multicast groups that hosts on the subnet serviced by an interface can join. By default, all groups are allowed on an interface. For <i>access-list-number</i> , specify an IP standard access list number. The range is 1 to 99.
Step 4	exit	Return to global configuration mode.
Step 5	access-list <i>access-list-number</i> { deny permit } <i>source</i> [<i>source-wildcard</i>]	Create a standard access list. - For <i>access-list-number</i>, specify the access list created in Step 3. - The deny keyword denies access if the conditions are matched. The permit keyword permits access if the conditions are matched. - For <i>source</i>, specify the multicast group that hosts on the subnet can join. (Optional) For <i>source-wildcard</i>, enter the wildcard bits in dotted decimal notation to be applied to the source. Place ones in the bit positions that you want to ignore. Recall that the access list is always terminated by an implicit deny statement for everything.
Step 6	end	Return to privileged EXEC mode.

	Command	Purpose
Step 7	show ip igmp 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 groups on an interface, use the **no ip igmp access-group** interface configuration command.

This example shows how to configure hosts attached to a port as able to join only group 255.2.2.2:

```
Switch(config)# access-list 1 255.2.2.2 0.0.0.0
Switch(config-if)# interface gigabitethernet1/0/1
Switch(config-if)# ip igmp access-group 1
```

Changing the IGMP Version

By default, the switch uses IGMP Version 2, which provides features such as the IGMP query timeout and the maximum query response time.

All systems on the subnet must support the same version. The switch does not automatically detect Version 1 systems and switch to Version 1. You can mix Version 1 and Version 2 hosts on the subnet because Version 2 routers or switches always work correctly with IGMPv1 hosts.

Configure the switch for Version 1 if your hosts do not support Version 2.

Beginning in privileged EXEC mode, follow these steps to change the IGMP version. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 3	ip igmp version {1 2}	Specify the IGMP version that the switch uses. Note If you change to Version 1, you cannot configure the ip igmp query-interval or the ip igmp query-max-response-time interface configuration commands.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip igmp interface [<i>interface-id</i>]	Verify your entries.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To return to the default setting, use the **no ip igmp version** interface configuration command.

Modifying the IGMP Host-Query Message Interval

The switch periodically sends IGMP host-query messages to discover which multicast groups are present on attached networks. These messages are sent to the all-hosts multicast group (224.0.0.1) with a time-to-live (TTL) of 1. The switch sends host-query messages to refresh its knowledge of memberships present on the network. If, after some number of queries, the software discovers that no local hosts are members of a multicast group, the software stops forwarding multicast packets to the local network from remote origins for that group and sends a prune message upstream toward the source.

The switch elects a PIM designated router (DR) for the LAN (subnet). The DR is the router or multilayer switch with the highest IP address for IGMPv2. For IGMPv1, the DR is elected according to the multicast routing protocol that runs on the LAN. The designated router is responsible for sending IGMP host-query messages to all hosts on the LAN. In sparse mode, the designated router also sends PIM register and PIM join messages toward the RP router.

Beginning in privileged EXEC mode, follow these steps to modify the host-query interval. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 3	ip igmp query-interval <i>seconds</i>	Configure the frequency at which the designated router sends IGMP host-query messages. By default, the designated router sends IGMP host-query messages every 60 seconds to keep the IGMP overhead very low on hosts and networks. The range is 1 to 65535.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip igmp interface [<i>interface-id</i>]	Verify your entries.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To return to the default setting, use the **no ip igmp query-interval** interface configuration command.

Changing the IGMP Query Timeout for IGMPv2

If you are using IGMPv2, you can specify the period of time before the switch takes over as the querier for the interface. By default, the switch waits twice the query interval controlled by the **ip igmp query-interval** interface configuration command. After that time, if the switch has received no queries, it becomes the querier.

You can configure the query interval by entering the **show ip igmp interface** *interface-id* privileged EXEC command.

Beginning in privileged EXEC mode, follow these steps to change the IGMP query timeout. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 3	ip igmp querier-timeout <i>seconds</i>	Specify the IGMP query timeout. The default is 60 seconds (twice the query interval). The range is 60 to 300.
Step 4	end	Return to privileged EXEC mode.

	Command	Purpose
Step 5	show ip igmp interface [<i>interface-id</i>]	Verify your entries.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To return to the default setting, use the **no ip igmp querier-timeout** interface configuration command.

Changing the Maximum Query Response Time for IGMPv2

If you are using IGMPv2, you can change the maximum query response time advertised in IGMP queries. The maximum query response time enables the switch to quickly detect that there are no more directly connected group members on a LAN. Decreasing the value enables the switch to prune groups faster.

Beginning in privileged EXEC mode, follow these steps to change the maximum query response time. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 3	ip igmp query-max-response-time <i>seconds</i>	Change the maximum query response time advertised in IGMP queries. The default is 10 seconds. The range is 1 to 25.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip igmp interface [<i>interface-id</i>]	Verify your entries.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To return to the default setting, use the **no ip igmp query-max-response-time** interface configuration command.

Configuring the Switch as a Statically Connected Member

Sometimes there is either no group member on a network segment or a host cannot report its group membership by using IGMP. However, you might want multicast traffic to go to that network segment. These are ways to pull multicast traffic down to a network segment:

- Use the **ip igmp join-group** interface configuration command. With this method, the switch accepts the multicast packets in addition to forwarding them. Accepting the multicast packets prevents the switch from fast switching.
- Use the **ip igmp static-group** interface configuration command. With this method, the switch does not accept the packets itself, but only forwards them. This method enables fast switching. The outgoing interface appears in the IGMP cache, but the switch itself is not a member, as evidenced by lack of an *L* (local) flag in the multicast route entry.

Beginning in privileged EXEC mode, follow these steps to configure the switch itself to be a statically connected member of a group (and enable fast switching). This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 3	ip igmp static-group <i>group-address</i>	Configure the switch as a statically connected member of a group. By default, this feature is disabled.
Step 4	end	Return to privileged EXEC mode.
Step 5	show ip igmp interface [<i>interface-id</i>]	Verify your entries.
Step 6	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To remove the switch as a member of the group, use the **no ip igmp static-group** *group-address* interface configuration command.

Configuring Optional Multicast Routing Features

These sections describe how to configure optional multicast routing features:

- Features for Layer 2 connectivity and MBONE multimedia conference session and set up:
 - [Enabling CGMP Server Support, page 40-32](#) (optional)
 - [Configuring sdr Listener Support, page 40-33](#) (optional)
- Features that control bandwidth utilization:
 - [Configuring an IP Multicast Boundary, page 40-35](#) (optional)

Enabling CGMP Server Support

The switch serves as a CGMP server for devices that do not support IGMP snooping but have CGMP client functionality. CGMP is a protocol used on Cisco routers and multilayer switches connected to Layer 2 Catalyst switches to perform tasks similar to those performed by IGMP. CGMP is necessary because the Layer 2 switch cannot distinguish between IP multicast data packets and IGMP report messages, which are both at the MAC-level and are addressed to the same group address.

Beginning in privileged EXEC mode, follow these steps to enable the CGMP server on the switch interface. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface that is connected to the Layer 2 Catalyst switch, and enter interface configuration mode.

	Command	Purpose
Step 3	ip cgmp [proxy]	Enable CGMP on the interface. By default, CGMP is disabled on all interfaces. Enabling CGMP triggers a CGMP join message. Enable CGMP only on Layer 3 interfaces connected to Layer 2 Catalyst switches. (Optional) When you enter the proxy keyword, the CGMP proxy function is enabled. The proxy router advertises the existence of non-CGMP-capable routers by sending a CGMP join message with the non-CGMP-capable router MAC address and a group address of 0000.0000.0000. Note To perform CGMP proxy, the switch must be the IGMP querier. If you configure the ip cgmp proxy command, you must manipulate the IP addresses so that the switch is the IGMP querier, which might be the highest or lowest IP address, depending on which version of IGMP is running on the network. An IGMP Version 2 querier is selected based on the lowest IP address on the interface. An IGMP Version 1 querier is selected based on the multicast routing protocol used on the interface.
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.
Step 7		Verify the Layer 2 Catalyst switch CGMP-client configuration. For more information, see the documentation that shipped with the product.

To disable CGMP on the interface, use the **no ip cgmp** interface configuration command.

When multiple Cisco CGMP-capable devices are connected to a switched network and the **ip cgmp proxy** command is needed, we recommend that all devices be configured with the same CGMP option and have precedence for becoming the IGMP querier over non-Cisco routers.

Configuring sdr Listener Support

The MBONE is the small subset of Internet routers and hosts that are interconnected and capable of forwarding IP multicast traffic. Other multimedia content is often broadcast over the MBONE. Before you can join a multimedia session, you need to know what multicast group address and port are being used for the session, when the session is going to be active, and what sort of applications (audio, video, and so forth) are required on your workstation. The MBONE Session Directory Version 2 (sdr) tool provides this information. This freeware application can be downloaded from several sites on the World Wide Web, one of which is <http://www.video.ja.net/mice/index.html>.

SDR is a multicast application that listens to a well-known multicast group address and port for Session Announcement Protocol (SAP) multicast packets from SAP clients, which announce their conference sessions. These SAP packets contain a session description, the time the session is active, its IP multicast group addresses, media format, contact person, and other information about the advertised multimedia session. The information in the SAP packet is displayed in the SDR Session Announcement window.

Enabling sdr Listener Support

By default, the switch does not listen to session directory advertisements.

Beginning in privileged EXEC mode, follow these steps to enable the switch to join the default session directory group (224.2.127.254) on the interface and listen to session directory advertisements. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be enabled for sdr, and enter interface configuration mode.
Step 3	ip sdr listen	Enable sdr listener support.
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 disable sdr support, use the **no ip sdr listen** interface configuration command.

Limiting How Long an sdr Cache Entry Exists

By default, entries are never deleted from the sdr cache. You can limit how long the entry remains active so that if a source stops advertising SAP information, old advertisements are not needlessly kept.

Beginning in privileged EXEC mode, follow these steps to limit how long an sdr cache entry stays active in the cache. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip sdr cache-timeout <i>minutes</i>	Limit how long an sdr cache entry stays active in the cache. By default, entries are never deleted from the cache. For <i>minutes</i> , the range is 1 to 4294967295.
Step 3	end	Return to privileged EXEC mode.
Step 4	show running-config	Verify your entries.
Step 5	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To return to the default setting, use the **no ip sdr cache-timeout** global configuration command. To delete the entire cache, use the **clear ip sdr** privileged EXEC command.

To display the session directory cache, use the **show ip sdr** privileged EXEC command.

Configuring an IP Multicast Boundary

Administratively-scoped boundaries can be used to limit the forwarding of multicast traffic outside of a domain or subdomain. This approach uses a special range of multicast addresses, called *administratively-scoped addresses*, as the boundary mechanism. If you configure an administratively-scoped boundary on a routed interface, multicast traffic whose multicast group addresses fall in this range can not enter or exit this interface, thereby providing a firewall for multicast traffic in this address range.

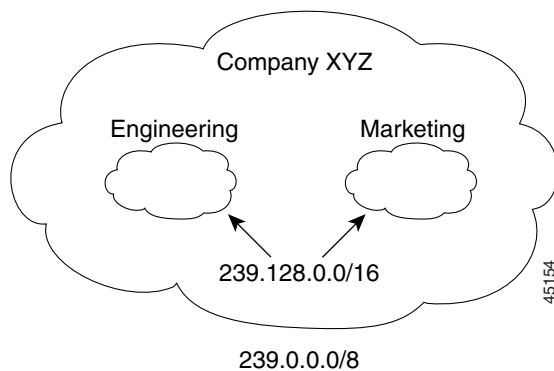


Note

Multicast boundaries and TTL thresholds control the scoping of multicast domains; however, TTL thresholds are not supported by the switch. You should use multicast boundaries instead of TTL thresholds to limit the forwarding of multicast traffic outside of a domain or a subdomain.

Figure 40-5 shows that Company XYZ has an administratively-scoped boundary set for the multicast address range 239.0.0.0/8 on all routed interfaces at the perimeter of its network. This boundary prevents any multicast traffic in the range 239.0.0.0 through 239.255.255.255 from entering or leaving the network. Similarly, the engineering and marketing departments have an administratively-scoped boundary of 239.128.0.0/16 around the perimeter of their networks. This boundary prevents multicast traffic in the range of 239.128.0.0 through 239.128.255.255 from entering or leaving their respective networks.

Figure 40-5 Administratively-Scoped Boundaries



You can define an administratively-scoped boundary on a routed interface for multicast group addresses. A standard access list defines the range of addresses affected. When a boundary is defined, no multicast data packets are allowed to flow across the boundary from either direction. The boundary allows the same multicast group address to be reused in different administrative domains.

The IANA has designated the multicast address range 239.0.0.0 to 239.255.255.255 as the administratively-scoped addresses. This range of addresses can then be reused in domains administered by different organizations. The addresses would be considered local, not globally unique.

Beginning in privileged EXEC mode, follow these steps to set up an administratively-scoped boundary. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	access-list <i>access-list-number</i> { deny permit } <i>source</i> [<i>source-wildcard</i>]	Create a standard access list, repeating the command as many times as necessary. - For <i>access-list-number</i>, the range is 1 to 99. - The deny keyword denies access if the conditions are matched. The permit keyword permits access if the conditions are matched. - For <i>source</i>, enter the number of the network or host from which the packet is being sent. (Optional) For <i>source-wildcard</i>, enter the wildcard bits in dotted decimal notation to be applied to the source. Place ones in the bit positions that you want to ignore. Recall that the access list is always terminated by an implicit deny statement for everything.
Step 3	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.
Step 4	ip multicast boundary <i>access-list-number</i>	Configure the boundary, specifying the access list you created in Step 2.
Step 5	end	Return to privileged EXEC mode.
Step 6	show running-config	Verify your entries.
Step 7	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To remove the boundary, use the **no ip multicast boundary** interface configuration command.

This example shows how to set up a boundary for all administratively-scoped addresses:

```
Switch(config)# access-list 1 deny 239.0.0.0 0.255.255.255
Switch(config)# access-list 1 permit 224.0.0.0 15.255.255.255
Switch(config)# interface gigabitethernet1/0/1
Switch(config-if)# ip multicast boundary 1
```

Configuring Basic DVMRP Interoperability Features

These sections contain this configuration information:

- [Configuring DVMRP Interoperability, page 40-37](#) (optional)
- [Configuring a DVMRP Tunnel, page 40-39](#) (optional)
- [Advertising Network 0.0.0.0 to DVMRP Neighbors, page 40-40](#) (optional)
- [Responding to mrinfo Requests, page 40-41](#) (optional)

For more advanced DVMRP features, see the “[Configuring Advanced DVMRP Interoperability Features](#)” section on page 40-41.

Configuring DVMRP Interoperability

Cisco multicast routers and multilayer switches using PIM can interoperate with non-Cisco multicast routers that use the DVMRP.

PIM devices dynamically discover DVMRP multicast routers on attached networks by listening to DVMR probe messages. When a DVMRP neighbor has been discovered, the PIM device periodically sends DVMRP report messages advertising the unicast sources reachable in the PIM domain. By default, directly connected subnets and networks are advertised. The device forwards multicast packets that have been forwarded by DVMRP routers and, in turn, forwards multicast packets to DVMRP routers.

You can configure an access list on the PIM routed interface connected to the MBONE to limit the number of unicast routes that are advertised in DVMRP route reports. Otherwise, all routes in the unicast routing table are advertised.



Note

The mroute protocol is a public-domain implementation of DVMRP. You must use mroute Version 3.8 (which implements a nonpruning version of DVMRP) when Cisco routers and multilayer switches are directly connected to DVMRP routers or interoperate with DVMRP routers over an MBONE tunnel. DVMRP advertisements produced by the Cisco IOS software can cause older versions of the mroute protocol to corrupt their routing tables and those of their neighbors.

You can configure what sources are advertised and what metrics are used by configuring the **ip dvmrp metric** interface configuration command. You can also direct all sources learned through a particular unicast routing process to be advertised into DVMRP.

Beginning in privileged EXEC mode, follow these steps to configure the sources that are advertised and the metrics that are used when DVMRP route-report messages are sent. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	access-list <i>access-list-number</i> { deny permit } <i>source</i> [<i>source-wildcard</i>]	Create a standard access list, repeating the command as many times as necessary. - For <i>access-list-number</i>, the range is 1 to 99. - The deny keyword denies access if the conditions are matched. The permit keyword permits access if the conditions are matched. - For <i>source</i>, enter the number of the network or host from which the packet is being sent. (Optional) For <i>source-wildcard</i>, enter the wildcard bits in dotted decimal notation to be applied to the source. Place ones in the bit positions that you want to ignore. Recall that the access list is always terminated by an implicit deny statement for everything.
Step 3	interface <i>interface-id</i>	Specify the interface connected to the MBONE and enabled for multicast routing, and enter interface configuration mode.

	Command	Purpose
Step 4	ip dvmrp metric <i>metric</i> [list <i>access-list-number</i>] [[<i>protocol process-id</i>] [dvmrp]]	Configure the metric associated with a set of destinations for DVMRP reports. - For <i>metric</i>, the range is 0 to 32. A value of 0 means that the route is not advertised. A value of 32 is equivalent to infinity (unreachable). (Optional) For list <i>access-list-number</i>, enter the access list number created in Step 2. If specified, only the multicast destinations that match the access list are reported with the configured metric. (Optional) For <i>protocol process-id</i>, enter the name of the unicast routing protocol, such as eigrp, igrp, ospf, rip, static, or dvmrp, and the process ID number of the routing protocol. If specified, only routes learned by the specified routing protocol are advertised in DVMRP report messages. (Optional) If specified, the dvmrp keyword allows routes from the DVMRP routing table to be advertised with the configured <i>metric</i> or filtered.
Step 5	end	Return to privileged EXEC mode.
Step 6	show running-config	Verify your entries.
Step 7	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To disable the metric or route map, use the **no ip dvmrp metric** *metric* [**list** *access-list-number*] [[*protocol process-id*] | [**dvmrp**]] or the **no ip dvmrp metric** *metric* **route-map** *map-name* interface configuration command.

A more sophisticated way to achieve the same results as the preceding command is to use a route map (**ip dvmrp metric** *metric* **route-map** *map-name* interface configuration command) instead of an access list. You subject unicast routes to route-map conditions before they are injected into DVMRP.

This example shows how to configure DVMRP interoperability when the PIM device and the DVMRP router are on the same network segment. In this example, access list 1 advertises the networks (198.92.35.0, 198.92.36.0, 198.92.37.0, 131.108.0.0, and 150.136.0.0) to the DVMRP router, and access list 2 prevents all other networks from being advertised (**ip dvmrp metric 0** interface configuration command).

```
Switch(config-if)# interface gigabitethernet1/0/1
Switch(config-if)# ip address 131.119.244.244 255.255.255.0
Switch(config-if)# ip pim dense-mode
Switch(config-if)# ip dvmrp metric 1 list 1
Switch(config-if)# ip dvmrp metric 0 list 2
Switch(config-if)# exit
Switch(config)# access-list 1 permit 198.92.35.0 0.0.0.255
Switch(config)# access-list 1 permit 198.92.36.0 0.0.0.255
Switch(config)# access-list 1 permit 198.92.37.0 0.0.0.255
Switch(config)# access-list 1 permit 131.108.0.0 0.0.255.255
Switch(config)# access-list 1 permit 150.136.0.0 0.0.255.255
Switch(config)# access-list 1 deny 0.0.0.0 255.255.255.255
Switch(config)# access-list 2 permit 0.0.0.0 255.255.255.255
```

Configuring a DVMRP Tunnel

The software supports DVMRP tunnels to the MBONE. You can configure a DVMRP tunnel on a router or multilayer switch if the other end is running DVMRP. The software then sends and receives multicast packets through the tunnel. This strategy enables a PIM domain to connect to the DVMRP router when all routers on the path do not support multicast routing. You cannot configure a DVMRP tunnel between two routers.

When a Cisco router or multilayer switch runs DVMRP through a tunnel, it advertises sources in DVMRP report messages, much as it does on real networks. The software also caches DVMRP report messages it receives and uses them in its RPF calculation. This behavior enables the software to forward multicast packets received through the tunnel.

When you configure a DVMRP tunnel, you should assign an IP address to a tunnel in these cases:

- To send IP packets through the tunnel
- To configure the software to perform DVMRP summarization

The software does not advertise subnets through the tunnel if the tunnel has a different network number from the subnet. In this case, the software advertises only the network number through the tunnel.

Beginning in privileged EXEC mode, follow these steps to configure a DVMRP tunnel. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	access-list <i>access-list-number</i> { deny permit } <i>source</i> [<i>source-wildcard</i>]	Create a standard access list, repeating the command as many times as necessary. • For <i>access-list-number</i>, the range is 1 to 99. • The deny keyword denies access if the conditions are matched. The permit keyword permits access if the conditions are matched. • For <i>source</i>, enter the number of the network or host from which the packet is being sent. • (Optional) For <i>source-wildcard</i>, enter the wildcard bits in dotted decimal notation to be applied to the source. Place ones in the bit positions that you want to ignore. Recall that the access list is always terminated by an implicit deny statement for everything.
Step 3	interface tunnel <i>number</i>	Specify a tunnel interface, and enter interface configuration mode.
Step 4	tunnel source <i>ip-address</i>	Specify the source address of the tunnel interface. Enter the IP address of the interface on the switch.
Step 5	tunnel destination <i>ip-address</i>	Specify the destination address of the tunnel interface. Enter the IP address of the mrouter.
Step 6	tunnel mode dvmrp	Configure the encapsulation mode for the tunnel to DVMRP.
Step 7	ip address <i>address mask</i> or ip unnumbered <i>type number</i>	Assign an IP address to the interface. or Configure the interface as unnumbered.
Step 8	ip pim [dense-mode sparse-mode]	Configure the PIM mode on the interface.

	Command	Purpose
Step 9	ip dvmrp accept-filter <i>access-list-number</i> [<i>distance</i>] neighbor-list <i>access-list-number</i>	Configure an acceptance filter for incoming DVMRP reports. By default, all destination reports are accepted with a distance of 0. Reports from all neighbors are accepted. - For <i>access-list-number</i>, specify the access list number created in Step 2. Any sources that match the access list are stored in the DVMRP routing table with distance. (Optional) For <i>distance</i>, enter the administrative distance to the destination. By default, the administrative distance for DVMRP routes is 0 and take precedence over unicast routing table routes. If you have two paths to a source, one through unicast routing (using PIM as the multicast routing protocol) and another using DVMRP, and if you want to use the PIM path, increase the administrative distance for DVMRP routes. The range is 1 to 255. - For neighbor-list <i>access-list-number</i>, enter the number of the neighbor list created in Step 2. DVMRP reports are accepted only by those neighbors on the list.
Step 10	end	Return to privileged EXEC mode.
Step 11	show running-config	Verify your entries.
Step 12	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To disable the filter, use the **no ip dvmrp accept-filter** *access-list-number* [*distance*] **neighbor-list** *access-list-number* interface configuration command.

This example shows how to configure a DVMRP tunnel. In this configuration, the IP address of the tunnel on the Cisco switch is assigned *unnumbered*, which causes the tunnel to appear to have the same IP address as port 1. The tunnel endpoint source address is 172.16.2.1, and the tunnel endpoint address of the remote DVMRP router to which the tunnel is connected is 192.168.1.10. Any packets sent through the tunnel are encapsulated in an outer IP header. The Cisco switch is configured to accept incoming DVMRP reports with a distance of 100 from 198.92.37.0 through 198.92.37.255.

```
Switch(config)# ip multicast-routing
Switch(config)# interface tunnel 0
Switch(config-if)# ip unnumbered gigabitethernet1/0/1
Switch(config-if)# ip pim dense-mode
Switch(config-if)# tunnel source gigabitethernet1/0/1
Switch(config-if)# tunnel destination 192.168.1.10
Switch(config-if)# tunnel mode dvmrp
Switch(config-if)# ip dvmrp accept-filter 1 100
Switch(config-if)# interface gigabitethernet1/0/1
Switch(config-if)# ip address 172.16.2.1 255.255.255.0
Switch(config-if)# ip pim dense-mode
Switch(config)# exit
Switch(config)# access-list 1 permit 198.92.37.0 0.0.0.255
```

Advertising Network 0.0.0.0 to DVMRP Neighbors

If your switch is a neighbor of an mrouter Version 3.6 device, you can configure the software to advertise network 0.0.0.0 (the default route) to the DVMRP neighbor. The DVMRP default route computes the RPF information for any multicast sources that do not match a more specific route.

Do not advertise the DVMRP default into the MBONE.

Beginning in privileged EXEC mode, follow these steps to advertise network 0.0.0.0 to DVMRP neighbors on an interface. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface that is connected to the DVMRP router, and enter interface configuration mode.
Step 3	ip dvmrp default-information {originate only}	Advertise network 0.0.0.0 to DVMRP neighbors. Use this command only when the switch is a neighbor of mrouted Version 3.6 machines. The keywords have these meanings: • originate—Specifies that other routes more specific than 0.0.0.0 can also be advertised. • only—Specifies that no DVMRP routes other than 0.0.0.0 are advertised.
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 prevent the default route advertisement, use the **no ip dvmrp default-information** interface configuration command.

Responding to mrinfo Requests

The software answers mrinfo requests sent by mrouted systems and Cisco routers and multilayer switches. The software returns information about neighbors through DVMRP tunnels and all the routed interfaces. This information includes the metric (always set to 1), the configured TTL threshold, the status of the interface, and various flags. You can also use the **mrinfo** privileged EXEC command to query the router or switch itself, as in this example:

```
Switch# mrinfo
171.69.214.27 (mm1-7kd.cisco.com) [version cisco 11.1] [flags: PMS]:
171.69.214.27 -> 171.69.214.26 (mm1-r7kb.cisco.com) [1/0/pim/querier]
171.69.214.27 -> 171.69.214.25 (mm1-45a.cisco.com) [1/0/pim/querier]
171.69.214.33 -> 171.69.214.34 (mm1-45c.cisco.com) [1/0/pim]
171.69.214.137 -> 0.0.0.0 [1/0/pim/querier/down/leaf]
171.69.214.203 -> 0.0.0.0 [1/0/pim/querier/down/leaf]
171.69.214.18 -> 171.69.214.20 (mm1-45e.cisco.com) [1/0/pim]
171.69.214.18 -> 171.69.214.19 (mm1-45c.cisco.com) [1/0/pim]
171.69.214.18 -> 171.69.214.17 (mm1-45a.cisco.com) [1/0/pim]
```

Configuring Advanced DVMRP Interoperability Features

Cisco routers and multilayer switches run PIM to forward multicast packets to receivers and receive multicast packets from senders. It is also possible to propagate DVMRP routes into and through a PIM cloud. PIM uses this information; however, Cisco routers and multilayer switches do not implement DVMRP to forward multicast packets.

These sections contain this configuration information:

- [Enabling DVMRP Unicast Routing, page 40-42](#) (optional)
- [Rejecting a DVMRP Nonpruning Neighbor, page 40-43](#) (optional)
- [Controlling Route Exchanges, page 40-45](#) (optional)

For information on basic DVMRP features, see the “[Configuring Basic DVMRP Interoperability Features](#)” section on page 40-36.

Enabling DVMRP Unicast Routing

Because multicast routing and unicast routing require separate topologies, PIM must follow the multicast topology to build loopless distribution trees. Using DVMRP unicast routing, Cisco routers, multilayer switches, and mrouted-based machines exchange DVMRP unicast routes, to which PIM can then reverse-path forward.

Cisco devices do not perform DVMRP multicast routing among each other, but they can exchange DVMRP routes. The DVMRP routes provide a multicast topology that might differ from the unicast topology. This enables PIM to run over the multicast topology, thereby enabling sparse-mode PIM over the MBONE topology.

When DVMRP unicast routing is enabled, the router or switch caches routes learned in DVMRP report messages in a DVMRP routing table. When PIM is running, these routes might be preferred over routes in the unicast routing table, enabling PIM to run on the MBONE topology when it is different from the unicast topology.

DVMRP unicast routing can run on all interfaces. For DVMRP tunnels, it uses DVMRP multicast routing. This feature does not enable DVMRP multicast routing among Cisco routers and multilayer switches. However, if there is a DVMRP-capable multicast router, the Cisco device can do PIM/DVMRP multicast routing.

Beginning in privileged EXEC mode, follow these steps to enable DVMRP unicast routing. This procedure is optional.

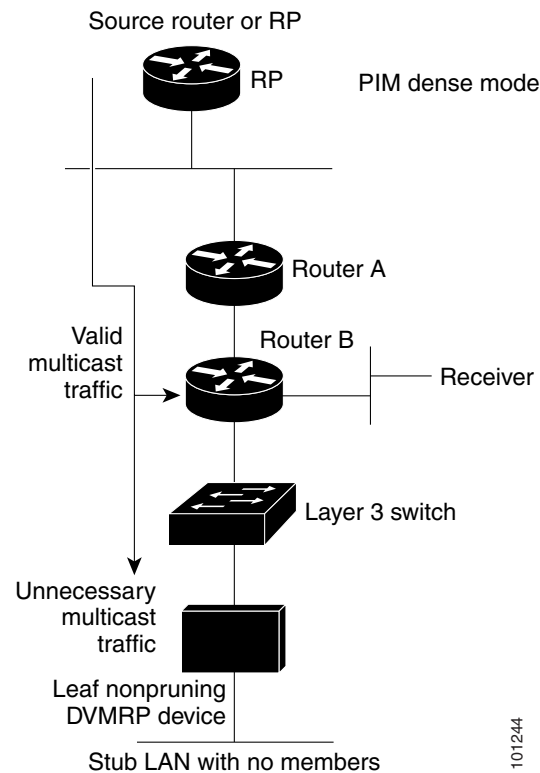
	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface that is connected to the DVMRP router, and enter interface configuration mode.
Step 3	ip dvmrp unicast-routing	Enable DVMRP unicast routing (to send and receive DVMRP routes). This feature is disabled by default.
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 disable this feature, use the **no ip dvmrp unicast-routing** interface configuration command.

Rejecting a DVMRP Nonpruning Neighbor

By default, Cisco devices accept all DVMRP neighbors as peers, regardless of their DVMRP capability. However, some non-Cisco devices run old versions of DVMRP that cannot prune, so they continuously receive forwarded packets, wasting bandwidth. [Figure 40-6](#) shows this scenario.

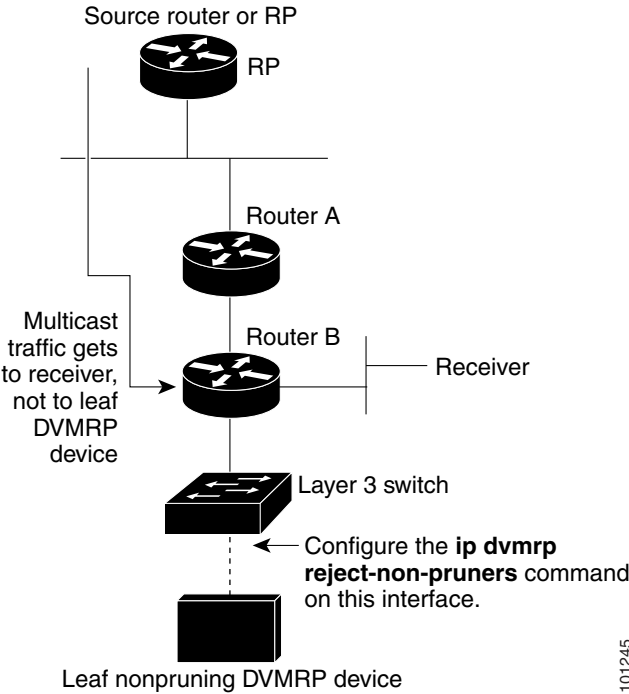
Figure 40-6 Leaf Nonpruning DVMRP Neighbor



You can prevent the switch from peering (communicating) with a DVMRP neighbor if that neighbor does not support DVMRP pruning or grafting. To do so, configure the switch (which is a neighbor to the leaf, nonpruning DVMRP machine) with the **ip dvmrp reject-non-pruners** interface configuration command on the interface connected to the nonpruning machine as shown in [Figure 40-7](#). In this case, when the switch receives DVMRP probe or report message without the prune-capable flag set, the switch logs a syslog message and discards the message.

101244

Figure 40-7 Router Rejects Nonpruning DVMRP Neighbor



Note that the **ip dvmrp reject-non-pruners** interface configuration command prevents peering with neighbors only. If there are any nonpruning routers multiple hops away (downstream toward potential receivers) that are not rejected, a nonpruning DVMRP network might still exist.

Beginning in privileged EXEC mode, follow these steps to prevent peering with nonpruning DVMRP neighbors. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface connected to the nonpruning DVMRP neighbor, and enter interface configuration mode.
Step 3	ip dvmrp reject-non-pruners	Prevent peering with nonpruning DVMRP neighbors.
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 disable this function, use the **no ip dvmrp reject-non-pruners** interface configuration command.

Controlling Route Exchanges

These sections describe how to tune the Cisco device advertisements of DVMRP routes:

- [Limiting the Number of DVMRP Routes Advertised, page 40-45](#) (optional)
- [Changing the DVMRP Route Threshold, page 40-45](#) (optional)
- [Configuring a DVMRP Summary Address, page 40-46](#) (optional)
- [Disabling DVMRP Autosummarization, page 40-48](#) (optional)
- [Adding a Metric Offset to the DVMRP Route, page 40-48](#) (optional)

Limiting the Number of DVMRP Routes Advertised

By default, only 7000 DVMRP routes are advertised over an interface enabled to run DVMRP (that is, a DVMRP tunnel, an interface where a DVMRP neighbor has been discovered, or an interface configured to run the **ip dvmrp unicast-routing** interface configuration command).

Beginning in privileged EXEC mode, follow these steps to change the DVMRP route limit. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip dvmrp route-limit <i>count</i>	Change the number of DVMRP routes advertised over an interface enabled for DVMRP. This command prevents misconfigured ip dvmrp metric interface configuration commands from causing massive route injection into the MBONE. By default, 7000 routes are advertised. The range is 0 to 4294967295.
Step 3	end	Return to privileged EXEC mode.
Step 4	show running-config	Verify your entries.
Step 5	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To configure no route limit, use the **no ip dvmrp route-limit** global configuration command.

Changing the DVMRP Route Threshold

By default, 10,000 DVMRP routes can be received per interface within a 1-minute interval. When that rate is exceeded, a syslog message is issued, warning that there might be a route surge occurring. The warning is typically used to quickly detect when devices have been misconfigured to inject a large number of routes into the MBONE.

Beginning in privileged EXEC mode, follow these steps to change the threshold number of routes that trigger the warning. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	ip dvmrp routehog-notification <i>route-count</i>	Configure the number of routes that trigger a syslog message. The default is 10,000 routes. The range is 1 to 4294967295.
Step 3	end	Return to privileged EXEC mode.
Step 4	show running-config	Verify your entries.
Step 5	copy running-config startup-config	(Optional) Save your entries in the configuration file.

To return to the default setting use the **no ip dvmrp routehog-notification** global configuration command.

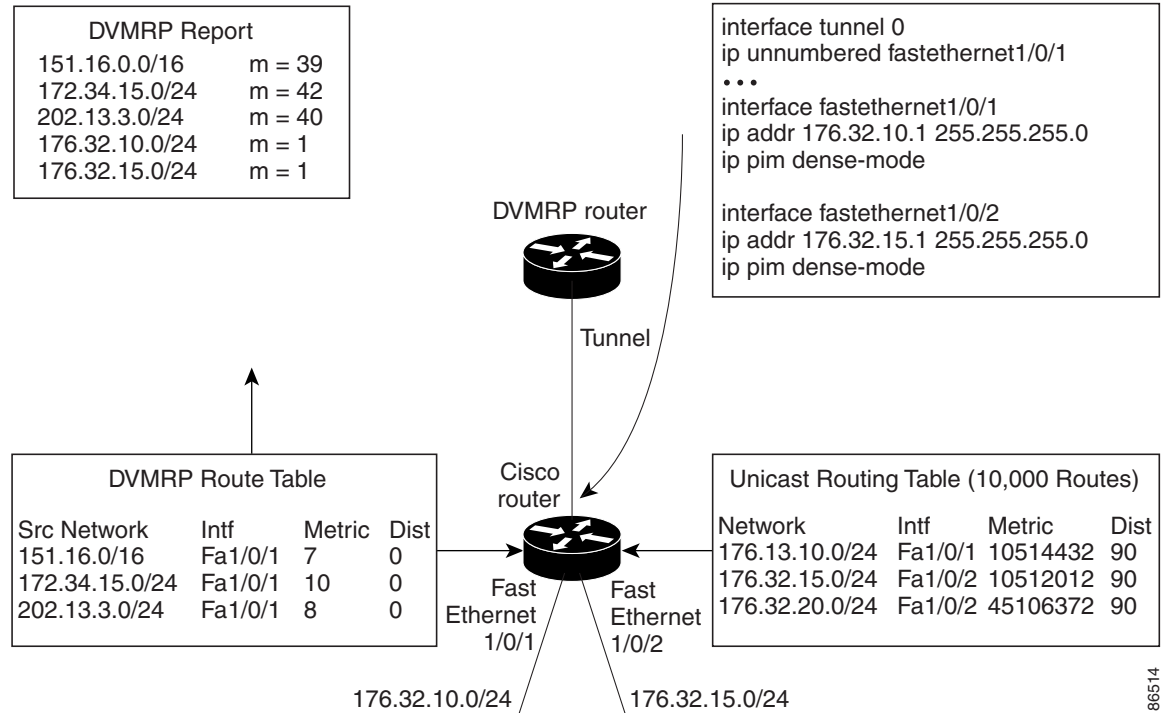
Use the **show ip igmp interface** privileged EXEC command to display a running count of routes. When the count is exceeded, ***** ALERT ***** is appended to the line.

Configuring a DVMRP Summary Address

By default, a Cisco device advertises in DVMRP route-report messages only connected unicast routes (that is, only routes to subnets that are directly connected to the router) from its unicast routing table. These routes undergo normal DVMRP classful route summarization. This process depends on whether the route being advertised is in the same classful network as the interface over which it is being advertised.

Figure 40-8 shows an example of the default behavior. This example shows that the DVMRP report sent by the Cisco router contains the three original routes received from the DVMRP router that have been poison-reversed by adding 32 to the DVMRP metric. Listed after these routes are two routes that are advertisements for the two directly connected networks (176.32.10.0/24 and 176.32.15.0/24) that were taken from the unicast routing table. Because the DVMRP tunnel shares the same IP address as Fast Ethernet port 1 and falls into the same Class B network as the two directly connected subnets, classful summarization of these routes was not performed. As a result, the DVMRP router is able to poison-reverse only these two routes to the directly connected subnets and is able to only RPF properly for multicast traffic sent by sources on these two Ethernet segments. Any other multicast source in the network behind the Cisco router that is not on these two Ethernet segments does not properly RPF-check on the DVMRP router and is discarded.

You can force the Cisco router to advertise the summary address (specified by the address and mask pair in the **ip dvmrp summary-address address mask** interface configuration command) in place of any route that falls in this address range. The summary address is sent in a DVMRP route report if the unicast routing table contains at least one route in this range; otherwise, the summary address is not advertised. In Figure 40-8, you configure the **ip dvmrp summary-address** command on the Cisco router tunnel interface. As a result, the Cisco router sends only a single summarized Class B advertisement for network 176.32.0.0.16 from the unicast routing table.

Figure 40-8 Only Connected Unicast Routes Are Advertised by Default

Beginning in privileged EXEC mode, follow these steps to customize the summarization of DVMRP routes if the default classful autosummarization does not suit your needs. This procedure is optional.

**Note**

At least one more-specific route must be present in the unicast routing table before a configured summary address is advertised.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface that is connected to the DVMRP router, and enter interface configuration command.
Step 3	ip dvmrp summary-address <i>address mask</i> [metric <i>value</i>]	Specify a DVMRP summary address. - For summary-address <i>address mask</i>, specify the summary IP address and mask that is advertised instead of the more specific route. (Optional) For metric <i>value</i>, specify the metric that is advertised with the summary address. The default is 1. The range is 1 to 32.
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 remove the summary address, use the **no ip dvmrp summary-address** *address mask* [**metric** *value*] interface configuration command.

Disabling DVMRP Autosummarization

By default, the software automatically performs some level of DVMRP summarization. Disable this function if you want to advertise all routes, not just a summary. In some special cases, you can use the neighboring DVMRP router with all subnet information to better control the flow of multicast traffic in the DVMRP network. One such case might occur if the PIM network is connected to the DVMRP cloud at several points and more specific (unsummarized) routes are being injected into the DVMRP network to advertise better paths to individual subnets inside the PIM cloud.

If you configure the **ip dvmrp summary-address** interface configuration command and did not configure **no ip dvmrp auto-summary**, you get both custom and autosummaries.

Beginning in privileged EXEC mode, follow these steps to disable DVMRP autosummarization. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface connected to the DVMRP router, and enter interface configuration mode.
Step 3	no ip dvmrp auto-summary	Disable DVMRP autosummarization.
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 re-enable auto summarization, use the **ip dvmrp auto-summary** interface configuration command.

Adding a Metric Offset to the DVMRP Route

By default, the switch increments by one the metric (hop count) of a DVMRP route advertised in incoming DVMRP reports. You can change the metric if you want to favor or not favor a certain route.

For example, a route is learned by multilayer switch A, and the same route is learned by multilayer switch B with a higher metric. If you want to use the path through switch B because it is a faster path, you can apply a metric offset to the route learned by switch A to make it larger than the metric learned by switch B, and you can choose the path through switch B.

Beginning in privileged EXEC mode, follow these steps to change the default metric. This procedure is optional.

	Command	Purpose
Step 1	configure terminal	Enter global configuration mode.
Step 2	interface <i>interface-id</i>	Specify the interface to be configured, and enter interface configuration mode.

	Command	Purpose
Step 3	ip dvmrp metric-offset [in out] <i>increment</i>	Change the metric added to DVMRP routes advertised in incoming reports. The keywords have these meanings: (Optional) in—Specifies that the increment value is added to incoming DVMRP reports and is reported in mrinfo replies. (Optional) out—Specifies that the increment value is added to outgoing DVMRP reports for routes from the DVMRP routing table. If neither in nor out is specified, in is the default. For <i>increment</i>, specify the value that is added to the metric of a DVMRP router advertised in a report message. The range is 1 to 31. If the ip dvmrp metric-offset command is not configured on an interface, the default increment value for incoming routes is 1, and the default for outgoing routes is 0.
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 setting, use the **no ip dvmrp metric-offset** interface configuration command.

Monitoring and Maintaining IP Multicast Routing

These sections describe how to monitor and maintain IP multicast routing:

- [Clearing Caches, Tables, and Databases, page 40-49](#)
- [Displaying System and Network Statistics, page 40-50](#)
- [Monitoring IP Multicast Routing, page 40-51](#)

Clearing Caches, Tables, and Databases

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

You can use any of the privileged EXEC commands in [Table 40-4](#) to clear IP multicast caches, tables, and databases:

Table 40-4 Commands for Clearing Caches, Tables, and Databases

Command	Purpose
clear ip cgmp	Clear all group entries the Catalyst switches have cached.
clear ip dvmrp route { * <i>route</i> }	Delete routes from the DVMRP routing table.

Table 40-4 Commands for Clearing Caches, Tables, and Databases (continued)

Command	Purpose
clear ip igmp group [<i>group-name</i> <i>group-address</i> <i>interface</i>]	Delete entries from the IGMP cache.
clear ip mroute [* <i>group</i> [<i>source</i>]]	Delete entries from the IP multicast routing table.
clear ip pim auto-rp <i>rp-address</i>	Clear the Auto-RP cache.
clear ip sdr [<i>group-address</i> “ <i>session-name</i> ”]	Delete the Session Directory Protocol Version 2 cache or an sdr cache entry.

Displaying System and Network Statistics

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



Note

This release does not support per-route statistics.

You can display information to learn 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.

You can use any of the privileged EXEC commands in [Table 40-5](#) to display various routing statistics:

Table 40-5 Commands for Displaying System and Network Statistics

Command	Purpose
ping [<i>group-name</i> <i>group-address</i>]	Send an ICMP Echo Request to a multicast group address.
show ip dvmrp route [<i>ip-address</i>]	Display the entries in the DVMRP routing table.
show ip igmp groups [<i>group-name</i> <i>group-address</i> <i>type number</i>]	Display the multicast groups that are directly connected to the switch and that were learned through IGMP.
show ip igmp interface [<i>type number</i>]	Display multicast-related information about an interface.
show ip mcache [<i>group</i> [<i>source</i>]]	Display the contents of the IP fast-switching cache.
show ip mpacket [<i>source-address</i> <i>name</i>] [<i>group-address</i> <i>name</i>] [detail]	Display the contents of the circular cache-header buffer.
show ip mroute [<i>group-name</i> <i>group-address</i>] [<i>source</i>] [summary] [count] [active kbps]	Display the contents of the IP multicast routing table.
show ip pim interface [<i>type number</i>] [count]	Display information about interfaces configured for PIM.
show ip pim neighbor [<i>type number</i>]	List the PIM neighbors discovered by the switch.
show ip pim rp [<i>group-name</i> <i>group-address</i>]	Display the RP routers associated with a sparse-mode multicast group.

Table 40-5 *Commands for Displaying System and Network Statistics (continued)*

Command	Purpose
show ip rpf { <i>source-address</i> <i>name</i> }	Display how the switch is doing Reverse-Path Forwarding (that is, from the unicast routing table, DVMRP routing table, or static mroutes).
show ip sdr [<i>group</i> “ <i>session-name</i> ” detail]	Display the Session Directory Protocol Version 2 cache.

Monitoring IP Multicast Routing

You can use the privileged EXEC commands in [Table 40-6](#) to monitor IP multicast routers, packets, and paths:

Table 40-6 *Commands for Monitoring IP Multicast Routing*

Command	Purpose
mrinfo [<i>hostname</i> <i>address</i>] [<i>source-address</i> <i>interface</i>]	Query a multicast router or multilayer switch about which neighboring multicast devices are peering with it.
mstat <i>source</i> [<i>destination</i>] [<i>group</i>]	Display IP multicast packet rate and loss information.
mtrace <i>source</i> [<i>destination</i>] [<i>group</i>]	Trace the path from a source to a destination branch for a multicast distribution tree for a given group.

