



DESIGN GUIDE

IPv6 MULTICAST DEPLOYMENT AND CONFIGURATION GUIDE

SCOPE OF THE DOCUMENT

This document covers the protocols and procedures necessary for the operation of IPv6 Multicast, and the Cisco implementation specifics and command references. The deployment scenarios for both intra- and interdomain multicast are provided.

An application note that includes basic configuration and applications for the end-to-end deployment is described in [1].

This document does not deal with Protocol Independent Multicast (PIM) and general multicast specific issues; familiarity with it is expected from IPv4 Multicast.

The first part deals with all available IPv6 group addressing formats and is followed by the multicast procedures.

This document is based on the information provided by Nidhi Bhaskar and Toerless Eckert and profited from the 6net IPv6 Multicast EFT testing led by Stig Venas (Uninet), Wim Biemolt (Surfnnet), Jerome Durand (Renater), Alex Gall, and Simon Leinen (Switch).

MULTICAST GROUP ADDRESSING

Terminology

Two main terms need to be distinguished in the IPv6 Multicast addressing:

- **Multicast group address**—The full 128 bit IPv6 address
- **Multicast group identifier**—Part of the full 128 bit address that is dedicated by the particular format definition as the specific group distinguisher

Generic Multicast Group Addresses Definition

General multicast address format is defined by RFC3513:



```
binary 11111111
```

at the start of the address identifies the address as being a multicast address.

flgs

$$+ - + - + - + - +$$

is a set of 4 flags: $|0|0|0|T|$

$$+ - + - + - + - +$$

This definition is represented by FF00::/8 in the prefix notation format.

The T flag set to 0 means permanent groups allocated by IANA (link local services—NTP servers and routing protocols, for example). The initial list of assigned permanent addresses is available in RFC2375. The most recent list of allocated addresses is available on the IANA site [\[2\]](#).

Groups used for the user transmissions must have the T flag always set to 1.

scop

This field is roughly the equivalent of IPv4 Scoped Zones. However, in IPv6 this concept is accomplished by using the **scop** bits in the IPv6 address to explicitly define the scope of the address as follows:

0	reserved
1	interface-local scope
2	link-local scope
3	reserved
4	admin-local scope
5	site-local scope
6	(unassigned)
7	(unassigned)
8	organization-local scope
9	(unassigned)
A	(unassigned)
B	(unassigned)
C	(unassigned)
D	(unassigned)
E	global scope
F	reserved

Multicast group traffic with scopes 1 and 2 must never be forwarded beyond the local link. Only traffic with scope 4 and higher can be forwarded across several router interfaces/links accordingly to the local definitions.

group-ID

Identifies the multicast group.

Unicast Prefix-Based Addresses

In IPv4, the problem of global multicast address allocation was partially solved using IPv4 GLOP (RFC3180) addressing, whereby the content provider would embed its Border Gateway Protocol (BGP) Autonomous System Number (ASN) in the middle two octets of a 233/8 multicast group address. This resulted in content providers having 256 group addresses that were uniquely assigned to them by virtue of their assigned ASNs.

RFC3306 defines an alternative way to IPv4 GLOP addressing, but instead of using embedded BGP ASN, the unicast prefix is used to identify this address range. It solves the problem of non-BGP users (not having any ASN) and gives some private multicast group range to everybody connected to the IPv6 Internet.

The format of a unicast-based multicast address is defined by RFC3306 as follows:

	8		4		4		8		8		64		32	
+-----+-----+-----+-----+-----+-----+-----+-----+														
	11111111		flgs		scop		reserved		plen		network prefix		group ID	
+-----+-----+-----+-----+-----+-----+-----+-----+														

binary 11111111

at the start of the address identifies the address as being a multicast address.

flgs

The RFC redefines the 4 bit “flag” field as follows:

```
+--+--+--+
| 0 | 0 | P | T |
+--+--+--+
```

- P = 0 indicates a multicast address that is not assigned based on the network prefix. This indicates a multicast address as defined in RFC3513.
- P = 1 indicates a multicast address that is assigned based on the network prefix.
- If P = 1, T MUST be set to 1; otherwise, the setting of the T bit is defined in [Generic Multicast Group Addresses Definition](#).

This address range is therefore represented as FF30::/12 in the prefix format.

plen

This field defines the number of significant bits in the “network prefix” field.

network prefix

This field contains the IPv6 unicast network prefix (left-justified) of the content provider. The maximum length of this field is 64 bits. Unused prefix bits must be 0.

group ID

This field contains the 32 bit IPv6 Multicast group ID that is assigned by the owner of the unicast prefix

Following are a few examples of the structure of unicast prefix-based multicast addresses:

- Global prefixes—A network with a unicast prefix of 3FFE:FFFF:1::/48 would also have a unicast prefix-based multicast prefix of FF3x:0030:3FFE:FFFF:0001::/96
- Source-Specific Multicast (SSM)—All IPv6 SSM multicast addresses will have the format FF3x::/96.

where ‘x’ is any valid scope.

SSM Address Range

The SSM [\[20\]](#) address range is a special case of the unicast-prefix based address defined by RFC3306 as follows:

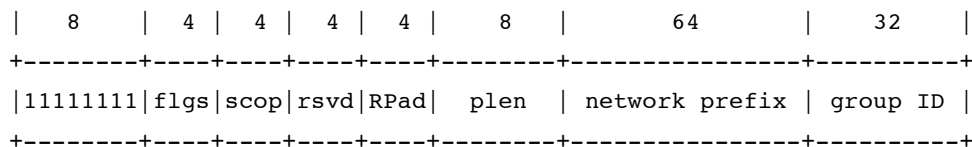
- Set P = 1.
- Set plen = 0.
- Set network prefix = 0.

These settings create a SSM range of FF3x::/96 (where ‘x’ is any valid scope value). Some specific group identifiers are excluded from the SSM range use by RFC3513. The source address field in the IPv6 header identifies the owner of the multicast address.

Embedded RP Addresses

RFC3956 [\[3\]](#) specifies a method to encode the Rendezvous Point (RP) address for a multicast group within a unicast prefix-based multicast address. This was accomplished by adding a new field for the RP address and by defining a new flag in the **flgs** field.

The format of an embedded RP address is as follows:

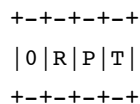


binary 11111111

at the start of the address identifies the address as being a multicast address.

flgs

The RFC redefines the 4 bit “flag” field as follows:



R = 1 indicates this is an embedded RP multicast address and contains the address of the PIM-SM RP. When R = 1, P MUST BE set to 1, and consequently T MUST also be set to 1, as specified in RFC3306 ([Generic Multicast Group Addresses Definition](#)).

This address range is therefore represented as FF70::/12 in the prefix format.

RPad

This field contains the RP address of the PIM-SM RP for this unicast prefix-based multicast address.

plen

This field defines the number of significant bits in the “network prefix” field.

network prefix

This field contains the IPv6 unicast network prefix (left justified) of the content provider. The maximum length of this field is 64 bits. Unused prefix bits do not have to be zeroed anymore—this RFC relaxes the strict requirement of RFC3306. The “plen” part of the prefix field is used to construct the RP address but the group address can contain a longer part of the unicast prefix.

group ID

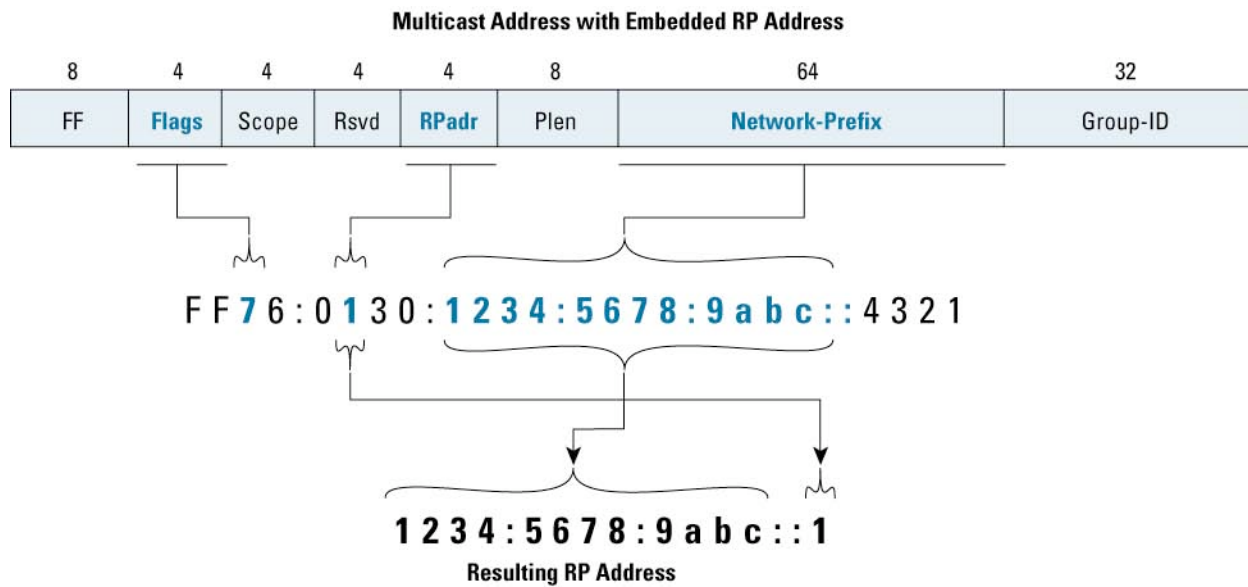
This field contains the 32 bit IPv6 Multicast group ID that is assigned by the owner of the unicast prefix.

Embedded RP multicast addresses are therefore a subset of the unicast prefix-based group address definition, and the RP address can be derived as follows:

1. Take the “plen” part of the “network prefix” field; add it with 0 bits to the full length of 128 bits
2. Replace the last 4 bits with the contents of “RPad”

[Figure 1](#) shows the mechanism for deriving the embedded RP address from the multicast group address in graphic form. The contents of the Network Prefix field are used as the prefix of the address, then the RPadr field is appended as the least significant bit of the RP address.

Figure 1. Constructing the Embedded RP Address



The IPv6 address representing the RP configured on the network devices needs to be allocated accordingly to these rules to match the embedded RP address.

Embedded RP Example (as provided by the IETF draft):

The network administrator of 2001:DB8::/32 wants to set up an RP for the network and all the customers. The administrator chooses network prefix=2001:DB8 and plen=32, and wants to use this addressing mechanism. The multicast addresses the administrator will be able to use are of the form:

FF7x:y20:2001:DB8:zzzz:zzzz:<group-id>

Where “x” is the multicast scope, “y” is the interface ID of the RP address, and “zzzz:zzzz” will be freely assignable to anyone. In this case, the address of the RP would be:

2001:DB8::y

and “y” could be anything from 1 to F, as 0 must not be used); the address 2001:DB8::y/128 is added on a router as a loopback address and injected to the routing system.

Multicast Address Mapping into MAC Address

The IPv6 Multicast address mapping into the Ethernet MAC address is defined by RFC2464 as follows:

An IPv6 packet with a multicast destination address DST, consisting of the sixteen octets DST[1] through DST[16], is transmitted to the Ethernet multicast address whose first two octets are the value 3333 hexadecimal and whose last four octets are the last four octets of DST (shown below).

Constructing the MAC Address

```
+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+
| 0 0 1 1 0 0 1 1 | 0 0 1 1 0 0 1 1 |
+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+
|   DST[ 13 ]      |   DST[ 14 ]      |
+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+
|   DST[ 15 ]      |   DST[ 16 ]      |
+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+
```

A multicast group address FF05:1::5 maps into a MAC address as shown below:

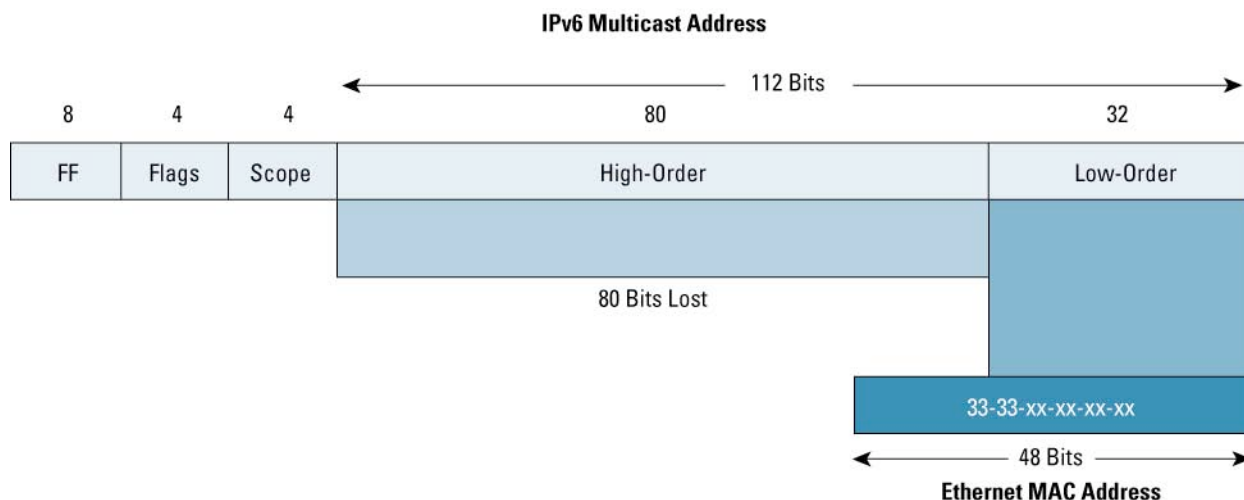
Mapping Group Address FF05:1::5 into the MAC Address

```
+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+
| 0 0 1 1 0 0 1 1 | 0 0 1 1 0 0 1 1 |
+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+
| 0 0 0 0 0 0 0 0 | 0 0 0 0 0 0 0 0 |
+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+
| 0 0 0 0 0 0 0 0 | 0 0 0 0 0 1 0 1 |
+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+--+
```

or 33:33:0:0:0:5 in the hexadecimal per byte notation.

[Figure 2](#) shows the mapping of the 128 bit Layer 3 IPv6 Multicast addresses to the 48 bit MAC address. Notice that there are 112 bits of significance to an IPv6 Multicast address (116 bits if the scope is included). However, only the lower 32 bits of the Layer 3 address is preserved, while 80 bits of information are lost during the mapping process. This means that it is possible for more than one IPv6 Multicast address to map to the same 48 bit MAC address.

Figure 2. MAC Address Mapping



In general, this overlap in Layer 3 to Layer 2 addressing *should* not be an issue—most applications should be using addresses that are unique in the lower 32 bits of address. However, when interdomain IPv6 Multicast is in use, there is a possibility that two globally scoped, network-prefix-based addresses could intersect in the lower 32 bits. This would result in both traffic streams being mapped into the same MAC address.

Cisco IOS Software Multicast Group Ranges

The interpretation of the multicast groups by Cisco IOS® Software can be checked using:

```
show ipv6 pim range
show ipv6 pim group-map
```

PIM Range

The PIM range command displays known group ranges to the router, as well as information about the RP for the particular ranges (if relevant):

```
7500> show ipv6 pim range
```

```
config SSM Exp: never Learnt from : ::
```

```
FF33::/32 Up: 04:44:07
FF34::/32 Up: 04:44:07
FF35::/32 Up: 04:44:07
FF36::/32 Up: 04:44:07
FF37::/32 Up: 04:44:07
FF38::/32 Up: 04:44:07
FF39::/32 Up: 04:44:07
FF3A::/32 Up: 04:44:07
FF3B::/32 Up: 04:44:07
FF3C::/32 Up: 04:44:07
FF3D::/32 Up: 04:44:07
FF3E::/32 Up: 04:44:07
FF3F::/32 Up: 04:44:07
```

```
config SM RP: 2001:efab:0:FE::1 Exp: never Learnt from : ::
  FF7B:140:2001:efab:0:FE::/96 Up: 04:44:06
config SM RP: 2001:abcd:14::2 Exp: never Learnt from : ::
  FF00::/8 Up: 04:44:07
```

- The first thirteen entries (FF33::/32—FF3F::/32) are the SSM ranges and are always hard-coded.
- The next entry (FF7B:140:2001:efab:0:FE::/96) is an example of an embedded RP entry. Embedded RP ranges (FF7x) appear only if there are groups active in that range (data appear, PIM, or MLD joins appear for the group range).
- Ordinary Sparse Mode group ranges appear together with the known RP. If only one RP is configured without an access list, only one entry for the whole of FF00::/8 range is displayed (although with implicit SSM range exclude).

PIM group-map

The PIM group-map commands displays the actual group to RP mapping for each known range with more information for every range and the RP (if relevant):

```
12000-1>show ipv6 pim group-map
```

SSM:

```
FF3E::/32*
```

```
RP      : ::
Protocol: SSM
Client  : config
Groups  : 0
Info    :
```

Sparse Mode:

```
FF00::/8*
```

```
RP      : 2001:abcd:14::2
Protocol: SM
Client  : config
Groups  : 2
Info    : RPF: Tu10,FE80::C0A8:F01
```

Embedded RP:

```
FF75:140:2001:efab:0:FE::/96*
```

```
RP      : 2001:efab:0:FE::1
Protocol: SM
Client  : Embedded
Groups  : 0
Info    : RPF: Tu11,FE80::C0A8:F02
```

MULTICAST COMPONENTS

IPv6 Multicast in Cisco IOS Software implements the new PIM-SM IETF draft [\[14\]](#), which introduces numerous protocol modifications for a clearer conceptual separation of all the functions necessary for the network to deliver multicast packets:

1. **Tree Information Base (TIB)**—Distribution tree topology information created by PIM and MLD procedures and kept by the router ([PIM TREE BUILDING PROCEDURES](#) and [USER TO NETWORK SIGNALING](#)). This information is stored internally by Cisco IOS Software on a Multicast Routing Protocol basis. (Since currently only the PIM protocol is implemented for IPv6 in Cisco IOS Software, this information appears as the “PIM Topology Table.”) Information from the topology table(s) is used to populate the Multicast Routing Information Base (MRIB). The relationship between topology tables and the MRIB is analogous to the unicast routing protocol databases (OSPF database or BGP table) and the Routing Information Base (RIB) in Cisco Express Forwarding.
2. **Multicast Routing Information Base (MRIB)**—Multicast routing information necessary to describe the multicast distribution trees (Shared Trees, Shortest-Path Trees, and Bidir Shared Trees, for example) that are used to forward multicast traffic. The content of the MRIB is a compilation of the information from the multicast topology table(s) produced by the different multicast routing protocols. (Since currently only the PIM protocol is implemented for IPv6 in Cisco IOS Software, only the “PIM Topology Table” exists and is used to create the MRIB. However, this separation of the topology tables and the MRIB provides the flexibility to implement future multicast routing protocols.)
3. **Multicast Forwarding Information Base (MFIB)**—Multicast forwarding information that is used to perform data forwarding ([MULTICAST DATA FORWARDING](#)). This table is a subset of the information in the MRIB along with the MAC header information to perform high-speed data forwarding of multicast traffic. The relationship between the MRIB and the MFIB is analogous to the unicast Routing Information Base (RIB) and the Forwarding Information Base (FIB) in Cisco Express Forwarding.

Unfortunately, there are some slight terminology differences between the acronyms used in the PIM specification and the acronyms used by Cisco IOS Software for the data structures above. These differences are detailed in Table 1.

Table 1. Nomenclature Differences Between Cisco IOS Software and the PIM Specification

PIM Spec		Cisco IOS Software	
MRIB	This table is similar to the normal unicast routing table (RIB) with the exception that it contains unicast routing information used to perform the multicast RPF check. This table permits incongruent unicast and multicast traffic routing.	?	Currently there is no equivalent data structure in Cisco IOS Software. Instead, Cisco IOS Software uses the unicast routing table (RIB) to perform the multicast RPF check.
TIB	Multicast topology information that describes the actual topology of the multicast distribution trees built by PIM.	TIB	This is the Cisco IOS Software multicast protocol topology table and is referred to as the “PIM Topology Table”.
Cisco IOS MRIB	This is considered by the PIM spec to be an implementation-specific data structure.	MRIB	Multicast routing information that describes the multicast distribution trees (Shared Trees, Shortest-Path Trees, Bidir Shared Trees) that are used to forward multicast traffic.
MFIB	This is considered by the PIM spec to be an implementation-specific data structure.	MFIB	Multicast forwarding information used to perform actual multicast data forwarding.

Enabling Multicast Routing

Similar to IPv4 Multicast, IPv6 Multicast data processing needs to be explicitly enabled in the global configuration mode using:

```
ipv6 multicast routing  
command.
```

UNICAST ROUTING

Because PIM is routing-protocol independent, it needs a separate unicast routing table to perform RPF checks. The table is created by the ordinary unicast routing protocols. In IPv6, this table can be created using:

Static Routing

The IPv4 unicast **ip route** command and its multicast alternative **ip mroute** command has been merged in IPv6 into one command:

```
ipv6 route <prefix> <next-hop> <admin-distance> unicast/multicast
```

By default, administrative distance is set to one and the route is used for both unicast forwarding and multicast RPF check. As soon as one of the unicast or multicast keywords is entered, it excludes the use of the route for the other option.

IS-IS for IPv6

IS-IS is specified in the ISO standard 10589 and has been designed to route primarily CLNS. The protocol data units are based on Type-Length-Value (TLV) message formats, which make the protocol easily extensible. RFC1195 extends IS-IS for routing of IPv4, [\[5\]](#) defines it for IPv6. Due the CLNS core tree, both IP protocols appear in the structure as leafs attached to the core tree. IS-IS therefore does not need to run separate Shortest Path First (SPF) calculations and separate databases for each of the IP protocols.

Configuring IS-IS for IPv6

IS-IS for IPv6 is available for Cisco IOS Software releases 12.0.21ST, 12.0.22S, 12.2.4T and later. The Cisco IOS Software syntax uses the format of address families if there is a need for IPv6-specific IS-IS commands/functions:

```
router isis
address-family ipv6
 redistribute static
 exit-address-family
```

Most of the configuration stays under the global **router isis** statement, so IPv6 IS-IS routing is enabled on the interfaces by the **ipv6 router isis** interface level command.

The full configuration documentation can be found at [\[6\]](#).

IS-IS in Multiprotocol Environment

In case IS-IS is used in a dual (or more) stack environment for routing IPv4, IPv6, and even CLNS on one infrastructure, RFC1195 specifies basic rules that must be kept in order for IS-IS to function properly—each routed protocol must be enabled fully on all links and routers in an entire Layer 1 area or whole of the Layer 2 backbone. The CLNS-based SPF tree calculation does not keep track of the routed protocol continuity; not keeping those rules can lead to black holing of traffic on links that were not enabled for the particular protocol.

Cisco IOS Software therefore (although not required by the specifications) enables by default the “protocols supported check”—it advertises in the hello messages the routed protocols enabled on the neighboring interfaces, and if the list of protocols does not fully match, the IS-IS adjacency is not closed. In some cases, especially on shared LAN segments, this check has to be disabled using:

```
router isis
 no adjacency-check
```

or

```
router isis
  address-family ipv6
    no adjacency-check
  exit-address-family
```

The applicability of those commands needs to be tested on a case-by-case basis, since the behavior can depend on the Cisco IOS Software release. Generally, non-IPv6-capable Cisco IOS Software can never close IS-IS adjacency to an IPv6-capable Cisco IOS Software release (since the rules to relax adjacency closure were introduced only with IPv6). The first command (right below `router isis`) resolves the case of a router with IPv4- and IPv6-enabled interface facing IPv6-only interface (the router in question will not seek an IPv4 address in the IS-IS hello coming from the IPv6-only interface). Similarly the IPv6 address family command resolves the case where the IPv4-only interface faces an IPv4- and IPv6-enabled interface (again, this time the IPv4 and IPv6 running router will not seek an IPv6 address in the coming hello and will close adjacency).

IS-IS Multitopology

IS-IS multitopology being developed in IETF [\[7\]](#) removes the limits on the supported protocols mentioned above (and provides the possibility to support several separated topologies for one protocol) at the expense of keeping separate databases and running separate SPF calculations for each topology. It also provides the possibility to create a separate IS-IS topology only for PIM RPF checks in a table not used for unicast forwarding. This feature is currently available in Cisco IOS Software Release 12.2.13T only for IPv6 unicast and can be enabled by:

```
router isis
  address-family ipv6 [unicast]
    multi-topology [transition]
```

Full configuration guide is available at [\[8\]](#).

This feature creates a separate IS-IS table on routers that support it. On routers that do not understand the new IS-IS messages, the corresponding IS-IS routes disappear from the IS-IS and routing table.

The new IS-IS table can be used for both unicast routing and multicast RPF checks.

RIP for IPv6

Routing Information Protocol (RIP) as originally defined is fully IPv4-specific. Its implementation for IPv6 routing represents a new protocol sometimes called RIPng. RFC2080 specifies the IPv6 alternative of RIPv2 for IPv4. Cisco IOS Software supports it since releases 12.0.21ST, 12.0.22S, and 12.2.4T. A RIPv6 routing process is enabled by:

```
ipv6 router rip
```

When configuring RIPng, each interface must be explicitly enabled for it; it is not possible to cover several interfaces (no `network` command available) under the **ipv6 router rip** command as in the case of the IPv4 alternative.

Full documentation is available at [\[9\]](#).

OSPF for IPv6

Similar to RIP, OSPF also required a new version specific only to IPv6—OSPFv3, which is described in RFC2740. OSPFv3 implements a separate OSPF database and a separate SPF calculation is run on this database. OSPFv3 has been supported in Cisco IOS Software since release 12.2.13T.

An OSPFv3 routing process is enabled by:

```
ipv6 router ospf
```

When configuring OSPFv3, each interface must be explicitly enabled for it. It is not possible to cover several interfaces under the **ipv6 router ospf** command as in the case of the IPv4 alternative.

Full documentation is available at [\[10\]](#).

Multitopology for OSPFv3 is currently only discussed in the IETF [\[11\]](#).

BGP for IPv6

The original implementation for IPv6 BGP in the tunnelled 6BONE experimental network was named as BGP4+ (the most common acronym for it is now MP-BGP or MBGP). It is based on the generic RFC2858 specifying the multiprotocol extensions for BGP—two new BGP attributes capable of transporting other than IPv4 prefixes in BGP. The BGP protocol format is possible to extend without any changes to the original specification in RFC1771 (the only IPv4-specific parameter) which must always be present in any BGP running device is a 32 bit-long BGP identifier.

RFC2545 defines encodings specific to IPv6, like the scope of BGP next-hop addresses, which should be advertised in the update messages. Address Family Identifiers (AFIs) are defined by IANA [\[12\]](#) and are 1 for IPv4 and 2 for IPv6. The Subsequent Address Family Identifier (SAFI) allows you to specify which prefix belongs to unicast routing table (SAFI=1) or multicast RPF table (SAFI=2).

Cisco IOS Software supports AFI=2 and SAFI=1 (IPv6 unicast) since releases 12.0.21ST, 12.0.22S, and 12.2.4T. The IPv6 Multicast address family will be available in releases 12.0.26S and 12.3T.

The configuration reference is available at [\[13\]](#).

Configuring BGP for Multicast Tables

The configuration tasks are identical to the unicast family:

```
address-family ipv6 multicast
neighbor 2001:abcd:0:C::1 activate
redistribute static
exit-address-family
```

Under the global **router bgp <ASN>** command, specify the IPv6 Multicast address family and enter the BGP commands that are relevant only to the BGP peers exchanging IPv6 SAFI multicast routes.

Originating Multicast Prefixes

There are several ways to originate BGP prefixes with multicast SAFI:

1. Network commands and redistribution—This needs to be configured under the multicast address family:

```
address-family ipv6 multicast
neighbor 2001:abcd:0:C::1 activate
network 2001::/16
```

```
redistribute static
exit-address-family
```

The network command will cause BGP prefix origination only if a unicast (but non-BGP) route matches exactly the configured prefix exists in the unicast routing table.

Similarly, the redistribute command example above will redistribute only static unicast routes into the multicast family (redistributing of routing protocols is also possible).

2. **Update translation from peers, which do not support multicast SAFI**—this configuration must appear under the unicast address family to the particular neighbor. Care needs to be taken to include unicast keywords in the command below—the intention usually will be to advertise both unicast and multicast (the multicast-only keyword would break unicast connectivity):

```
address-family ipv6
neighbor 2001:abcd:0:D::2 translate-update ipv6 multicast unicast
neighbor 2001:abcd:0:D::2 activate
no synchronization
exit-address-family

address-family ipv6 multicast
neighbor 2001:abcd:0:D::2 activate
```

To enable the update translation to work, the peer needs to be activated under the address-family IPv6 Multicast although it does not support it (a more detailed explanation of the use of this configuration is provided in [Translate Updates from Non-MP-BGP Peer](#)).

MP-BGP Incompatibilities

The set of BGP RFCs mentioned above (especially RFC1771, RFC2858, and RFC3392) is not specific enough to make all BGP implementations smoothly interoperate:

1. **BGP next-hop handling**—RFC1771 mandates the presence of the IPv4 NEXT_HOP attribute in every BGP UPDATE message. Obviously, this attribute does not have any sense in pure IPv6 or IPv4 Multicast environments, as the multiprotocol attribute carries its own next-hop value. (**Note:** Cisco BGP is strictly compliant to RFC1771 and will reset the BGP session if it does not see the original NEXT_HOP attribute. This issue is resolved in the new BGP and MP-BGP drafts.)
2. **Capabilities advertisements**—The capability to handle other than IPv4 address families is advertised during the BGP session setup in the options field of the BGP OPEN message. If one BGP router has not only multiprotocol neighbors, but also IPv4 unicast only (RFC1771 only) neighbors, the MP-BGP-capable code will encode IPv4 in the multiprotocol attributes and will require to see from the IPv4-only peer the multiprotocol capability advertisement. This is compliant with RFC3392, but prevents establishing a BGP session with those “older” BGP peers due to the capability mismatch—no common capability advertised (although both of them can handle IPv4 unicast).

IPv6 RPF Check

All the routing protocols mentioned so far can contribute to a unicast routing table in a router. Similar to unicast route selection from several possible alternatives, multicast RPF check needs to choose one from all available routes. There are only few rules for IPv6 RPF check:

1. Exclude IPv6 BGP unicast (SAFI=1) routes by default. It can be enabled by the **ipv6 rpf use-bgp** command.
2. Do longest match on the multicast packet source address across all available routing protocols tables in the exactly same way the lookup for unicast packet forwarding is done. Include all multicast-specific tables (currently only MP-BGP and static multicast routes). There is no configuration option to change this behavior.
3. Break ties using the administrative distances if there are prefixes of the same length in several protocol tables. If the administrative distance is still the same, prefer static routes above MP-BGP and MP-BGP above unicast routing protocols tables.

Recursive Route Lookups

In the BGP environment, the RPF lookup will first return the BGP next-hop IP address, which is typically not directly connected. Next resolution needs to be done to find a directly connected next-hop. A similar situation can occur when static routes are in use.

Every recursive lookup for the next-hops of a unicast BGP (or other unicast route that needs recursion to resolve) route is performed only in the unicast routing table. Recursive lookup for the next-hops of MP-BGP routes (or any recursive multicast route) will use not only the unicast table, but also MP-BGP and any other multicast routing tables. The maximum recursion depth is 8.

RPF Interface and Neighbor Selection

The new PIM-SM draft [\[14\]](#) defines explicitly the RPF interface and RPF neighbor and when they are used. Cisco IOS Software is compliant with these definitions.

RPF Interface

The RPF interface is the interface chosen by unicast routing to reach the multicast source or the RP (the presence of a PIM neighbor is not required and even the PIM configuration on the interface—PIM enabled/disabled—is not checked). The RPF interface is also used to validate that multicast data is arriving from the correct interface.

RPF Neighbor

The RPF neighbor is defined as the PIM neighbor residing on the RPF interface toward the multicast source up the Shortest-Path Tree (or the RP in the case of a Shared Tree). PIM Join/Prune message can be sent only to the RPF neighbor, the existence of just RPF interface selected by unicast routing (but possibly without a PIM neighbor on that interface) is not enough.

See [Monitoring Unicast Routing Tables](#) for more details and differences between the two terms defined here.

Equal-Cost Multipath Handling

Equal-cost paths to a multicast source are defined as several possible paths available in the same routing protocol table (several paths available in the OSPF table only, for example), which have equal metric to reach the source.

IPv6 RPF check always considers all possible paths in the RPF interface selection and performs a hash function using the last 32 bits of the source address and the last 32 bits of the next-hop address to calculate the next-hop priority. The interface with the highest priority next-hop is selected as the RPF interface. This design decision is compliant with RFC2991, which does not suggest/support the selection of just one link from a set of equal cost paths. The check for a PIM neighbor on the RPF interface is not performed.

Monitoring Unicast Routing Tables

The content of the unicast routing protocol tables can be checked using the ordinary routing-protocol-specific commands. In addition, the RPF information can be found using:

```
7500>show ipv6 rpf 2001:abcd:0:B::1
RPF information for 2001:abcd:0:B::1
  RPF interface: POS4/1/0
  RPF neighbor: FE80::208:E2FF:FE3C:300
  RPF route/mask: 2001:abcd:0:B::/64
  RPF type: MBGP
  RPF recursion count: 0
```

Metric preference: 20

Metric: 20

Note: The term “RPF neighbor” here means only the link local IP address of the routing protocol neighbor found on the interface; it is not an RPF neighbor in the sense of the definition in [RPF Neighbor](#).

The example output below shows the significance of RPF interface and neighbor. The RPF interface points to an interface that does not have any PIM neighbor (POS4/1/0 does not appear in the output below); therefore, the RPF neighbor does not exist here and PIM can not forward any Join/Prune messages over this interface:

```
7500>show ipv6 pim neighbor
```

Neighbor Address	Interface	Uptime	Expires	DR	pri	Bidir
FE80::208:A4FF:FEA7:A406	FastEthernet4/0/1	6d04h	00:01:37	0		B

Currently the unicast routing protocol (if other than BGP) name is not displayed (as it is in the case of IPv4 RPF), but only the corresponding administrative distance of the chosen protocol entry:

```
12000-1>show ipv6 rpf 2001:abcd:14::2
```

RPF information for 2001:abcd:14::2

RPF interface: POS1/0

RPF neighbor: FE80::208:E2FF:FE3B:EA00

RPF route/mask: 2001:abcd:14::2/128

RPF type: Unicast

RPF recursion count: 0

Metric preference: 115

Metric: 410

PIM TREE BUILDING PROCEDURES

PIMv6

PIM is specified in RFC2362, and there is a newly developed IETF draft [\[14\]](#). Both specifications are address-family independent; all addresses are encoded in the format of the “encoded-unicast address.” An encoded-unicast address takes the following format:

0	1	2	3
0 1 2 3 4 5 6 7 8 9 0 1 2 3 4 5 6 7 8 9 0 1 2 3 4 5 6 7 8 9 0 1			
+--+			
Addr Family		Encoding Type unicast Address	
+--+...			

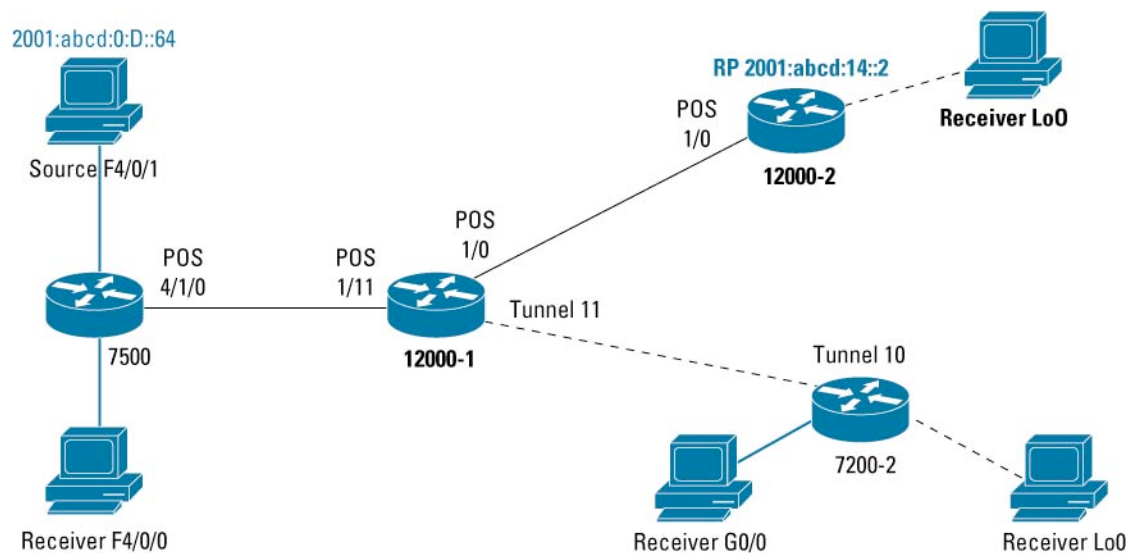
where the address family is the same number used in MP-BGP and assigned by IANA [\[12\]](#). This allows PIM implementation for IPv6.

Cisco IOS Software implements draft [\[14\]](#) for PIMv6, which brings changes in the commands outputs and population of the PIM states as described in the next sections.

PIM Topology Example

The outputs shown in this section are taken from the network on [Figure 3](#). The receivers on physical interfaces are simulated using static MLD join ([Configuring and Monitoring MLD](#)); receivers on Loopback are simulated using the MLD join-group command, which causes the router itself to join and receive the data.

Figure 3. Example Multicast Network



Configuring Interfaces for PIMv6

When IPv6 Multicast is enabled globally with the **ipv6 multicast-routing** command as discussed in [Enabling Multicast Routing](#), all IPv6-enabled interfaces on the router are automatically enabled also for PIMv6; no other configuration is necessary. Since this is the default configuration of each interface, the **ipv6 pim** command is assumed and does not appear in the interface configuration.

To disable PIM on a particular interface, use the:

```
no ipv6 pim
```

interface-level command. When PIM is disabled on the interface, multicast data forwarding over the interface is disabled.

The PIM configuration of all interfaces can be checked using the following command:

```
12000-2>show ipv6 pim interface
```

Interface	PIM	Nbr	Hello	DR
		Count	Intvl	Prior
Ethernet0	on	0	30	1
Address: FE80::208:E2FF:FE3C:3FF				
DR	: this system			

```
POS1/0          on 1 30 1
Address: FE80::208:E2FF:FE3C:300
DR      : this system
```

It is important to note that this behavior is in contrast to the default behavior of Cisco IOS Software in IPv4, where it is necessary to specifically enable PIMv4 on each interface using either **ip pim sparse-mode**, **ip pim dense-mode**, or **ip pim sparse-dense-mode**. In addition, the Cisco IPv6 implementation has eliminated these commands along with their respective operating modes. Instead, when PIMv6 is enabled on an interface, the interface *always* operates in sparse mode (dense mode operation is not currently implemented). This eliminates the complexity of interface mode management, which has historically been a source of confusion in IPv4. More details on the differences in the IPv4 and IPv6 implementations can be found in [Appendix A](#).

Group Modes

In Cisco IOS Software, each IP Multicast group operates in exactly one mode of PIM. This group mode is determined by configuration or learned via dynamic protocols described below, details on the differences in the IPv4 and IPv6 implementations can be found in [Appendix A](#).

General Any Source Multicast (ASM)

The group ranges specified in [Generic Multicast Group Addresses Definition](#) and [Unicast Prefix-Based Addresses](#) will, by default, always be handled as sparse mode groups. If Cisco IOS Software does not know an RP for a group, then a default group mode is used (sparse mode for IPv6, for example).

SSM

The SSM group range ([SSM Address Range](#)) is hard-coded and will always be handled in the SSM mode (no wildcard joins processed, no RP definition for the group range required).

Embedded RP Groups

Embedded group ranges ([Embedded RP Addresses](#)) cannot be predefined (the RP address is not preliminarily known) in the SSM sense, but the router interprets the embedded RP group address only when first PIM joins for that group appear or first data appear from the directly connected source—before that, it is not possible to determine any RP for the group.

The group mode can be displayed using the:

```
show ipv6 pim range-list
```

command ([PIM Range](#) for an output example).

By default, when multicast routing is enabled, the group modes are as follows [\[16\]](#):

ffX[0-2]::/16	Non-Routable.
ff3X::/32	SSM
ff00::/8	None

PIM Neighbor Address Resolution

In order for PIM to operate properly, it needs to correctly resolve the addresses of its PIM neighbors. This is essential for propagation of PIM joins—the unicast routing table will provide the information about the RPF interface toward the multicast data source or toward the RP, but the join will only be forwarded if a PIM neighbor is also found on that interface. The unicast routing process in the IPv6 case has several choices of what neighbor IPv6 address to return—link local, global, or any other type configured on the interface. PIM hellos and PIM neighbor discovery is done using link local addresses only. This may cause conflict and the PIM neighbor not being correctly resolved, even if there is one. The easiest example is when configuring a static multicast route using a global address:

```
ipv6 route 2001::/16 Tunnel10 2002::1 multicast
```

The RPF check lookup will return the global address:

```
7200-2>show ipv6 rpf 2001:abcd:14::2
```

RPF information for 2001:abcd:14::2

```
RPF interface: Tunnel10
RPF neighbor: 2002::1
RPF route/mask: 2001::/16
RPF type: Mroute
RPF recursion count: 0
Metric preference: 1
Metric: 0
```

PIM messages are sent using only link local addresses as the source IPv6 address for the PIM packets. A PIM neighbor will therefore be identified using only link local addresses, while the RPF lookup in the example above is returning a global IPv6 address of the neighbor:

```
7200-2>show ipv6 pim neighbor
```

Neighbor Address	Interface	Uptime	Expires	DR	pri	Bidir
FE80::C0A8:F01	Tunnel10	06:45:07	00:01:21	1		B

This problem has been resolved by adding a PIM option to the PIM hellos used to advertise all the IPv6 addresses on the particular PIM interface. This information is used to find a match between the unicast next-hop IPv6 address and a PIM neighbor IPv6 address. The initial IETF draft for this option is “*draft-suz-pim-upstream-detection*” [\[15\]](#), later on it became part of the PIM specification [\[14\]](#).

It can be checked in Cisco IOS Software using:

```
7200-2>show ipv6 pim neighbor detail
```

Neighbor Address(es)	Interface	Uptime	Expires	DR	pri	Bidir
FE80::C0A8:F01	Tunnel10	06:50:56	00:01:30	1		B
2002::1						

If more than the link local IPv6 address appears for each neighbor, this PIM option is implemented.

The PIM hello option type is not yet allocated by IANA: <http://www.iana.org/assignments/pim-hello-options>

Cisco uses 65001 while Juniper uses 24, which can cause interoperability issues.

Registering

The PIM-SM Draft [14] suggests that source registering be accomplished using a virtual tunnel interface. This use of virtual tunnel interfaces permits consistent PIM state handling for the registration process. During the registration process, the tunnel interface appears like any other interface in the Outgoing Interface List for the multicast data source (S,G) state with all the rules valid for the state management. In Cisco IOS Software implementations, an automatic tunnel is created as soon as an RP is known; one virtual tunnel for each active RP in the network. While the PIM-SM Draft [14] suggests that the tunnel should be deleted after each process of registering, Cisco IOS Software keeps each tunnel as long as the RP is known. The additional implementation-specific advantage of these tunnel interfaces is simplification of the register data encapsulation—it does not have to be handled specifically in the PIM part of the code. Instead, generic IP code can be used to perform the encapsulation such that the PIM Register packet is just forwarded into the tunnel for encapsulation. Use of generic tunnelling code in Cisco IOS Software enables the possible handling of PIM Register packets in fast (not process-switched) path if available.

These virtual tunnels are always unidirectional (transmit only) and automatic—the tunnel interface status immediately goes to **up** when it is created. However, the line protocol stays in **down** status until there is a valid RPF interface to the RP (for example, unicast connectivity through unicast BGP in the default configuration is not enough, as BGP is not used for RPF check) and also a unicast route exists in the unicast RIB to the RP. Sources can successfully register only when the tunnel interface is fully up.

It is important to note that while all PIM Register messages from the registering routers (first-hop routers) are sent to the RP via these virtual tunnels, all PIM Register-Stop messages are sent directly from the RP to the registering router and do not use virtual tunnels.

The handling of dynamic changes of RP information is not fully resolved in the first IPv6 Multicast implementation—it is a generic Cisco IOS Software issue, which can not handle properly deleting of interfaces (the register tunnels in this case) and reusing of the same interface number by a newly created tunnel. This can cause problems when BSR and embedded RP are used to distribute RP information (when the RP information dynamically changes).

The output of the following command displays the existing RP Register tunnels at the non-RP routers in the network:

```
Router> show ipv6 pim tunnel
```

```
Tunnel2*
```

```
Type   : PIM Encap  
RP      : 2001:abcd:25::1  
Source: 2001:abcd:0:3::1
```

```
Tunnel1*
```

```
Type   : PIM Encap  
RP      : 2001:efab:0:FE::1  
Source: 2001:abcd:0:C::1
```

In the example above, notice that there are two active RPs in the network for two different group ranges. **Tunnel1** connects to RP 2001:efab:0:FE::1 and **Tunnel2** connects to RP 2001:abcd:25::1.

The RP tunnels and the registering traffic can be monitored using ordinary interface commands:

```
7500>show interface tunnel 1
```

Tunnell is up, line protocol is up

Tunnel source 2001:abcd:0:C::2 (POS4/1/0), destination 2001:abcd:14::2

Tunnel protocol/transport PIM/IPv6, key disabled, sequencing disabled

Checksumming of packets disabled

Tunnel is transmit only

Last input never, output 00:00:00, output hang never

Last clearing of "show interface" counters 01:04:17

Input queue: 0/75/0/0 (size/max/drops/flushes); Total output drops: 0

Queueing strategy: fifo

Output queue: 0/0 (size/max)

5 minute input rate 0 bits/sec, 0 packets/sec

5 minute output rate 93000 bits/sec, 39 packets/sec

The PIM protocol traffic (including registering) can be displayed using:

```
7500>show ipv6 pim traffic
```

PIM Traffic Counters

Elapsed time since counters cleared: 4w4d

	Received	Sent
Valid PIM Packets	545180	326944
Hello	232748	355759
Join-Prune	15563	41200
Register	0	24265
Register Stop	24279	0
Assert	0	0
Bidir DF Election	0	0

Errors:

Malformed Packets	0
Bad Checksums	0
Send Errors	3
Packet Sent on Loopback Errors	94277
Packets Received on PIM-disabled Interface	0
Packets Received with Unknown PIM Version	0

When registering, the tunnel appears in the outgoing interface list. The view of the Cisco 7500 Series router ([Figure 3](#)), which is the first hop router for the source, is:

```
7500> show ipv6 pim topology

(2001:abcd:0:D::64,FF05:1::5)
SM SPT UP: 6d05h JP: Join(never) Flags: KAT(00:01:02) RA SR
RPF: FastEthernet4/0/1,2001:abcd:0:D::64*
  Tunnel0          5d02h      fwd
```

The above output reflects the status of the source registration before the RP has sent back any (S,G) Joins or Register-Stop messages. (**Note:** In order to catch this transient event, the registering source IP address was made temporarily unreachable at the RP.)

The PIM topology command displays the following flags when a source is registering:

KAT—Keep Alive Timer—A timer associated with each locally connected source or with each source at the RP or turnaround router. The source is assumed to be alive for 210 seconds; after this, the timer elapses and the router sends probes internally (an internal check of the packet counters if any data was sent recently) to test that the source is alive. If successful, the router resets the timer to 210 seconds; if not, the entry is deleted after 65s.

RA—Really Alive—Signifies a source that has already a (S,G) state created and the data has been received from the source.

SR—Sending Registers—Appears only temporarily when data registers are sent.

```
7500> show ipv6 mroute

(2001:abcd:0:D::64, FF05:1::5), 6d05h/00:00:04, flags: SFJT
  Incoming interface: FastEthernet4/0/1
  RPF nbr: 2001:abcd:0:D::64, Registering
  Outgoing interface list:
    Tunnel0, Forward, 5d02h/never
```

The meaning of flags in the **show ipv6 mroute** command is the same as for the IPv4 outputs.

On the RP, at least two tunnels are always created as shown in the example below:

```
12000-2> show ipv6 pim tunnel
```

```
Tunnel0*
  Type   : PIM Encap
  RP     : 2001:abcd:14::2
  Source : 2001:abcd:14::2
Tunnel1*
  Type   : PIM Decap
  RP     : 2001:abcd:14::2
  Source : -
```

The first tunnel at the RP (**Tunnel0** in the example above) is a transmit-only virtual tunnel that is used for the registering of any sources locally connected to the RP. The second tunnel at the RP (**Tunnel1** in the example above) is a receive-only virtual tunnel used for decapsulating all incoming registers from all remote designated routers. There is not a one-to-one correspondence of virtual tunnels at the RP to the virtual tunnels originating at the designated routers in the network. Instead, only a single **PIM Decap** tunnel is used at the RP to decapsulate *all* PIM Register traffic arriving from *all* designated routers.

The RP state (taken from the 12000-2 router on [Figure 3](#)) for the active source is exactly the same (including the flags) during registering and native data delivery:

```
12000-2> show ipv6 pim topology
```

```
(2001:abcd:0:D::64,FF05:1::5)
```

```
SM SPT UP: 00:29:37 JP: Join(00:00:44) Flags: KAT(00:02:59) RA RR
```

```
RPF: POS1/0,FE80::208:E2FF:FE3C:300
```

```
No interfaces in immediate olist
```

The only different flag as compared with source first-hop Cisco 7500 router is the RR flag set instead of SR:

RR—Register Received—Determines that the source has registered to this RP.

Restricting PIM Registers at the RP

The incoming PIM Register packets can be restricted at the RP using:

```
ipv6 pim accept-register
```

command with either:

```
list <acl_name> option to allow only particular sources to register
```

or

```
route-map <rmap_name> option to filter sources based on more criteria.
```

Distributing RP Information

The following mechanisms can be used in IPv6 Multicast to configure the RP information on the network devices:

Static Configuration

The RP can be configured statically on each network device using the command:

```
ipv6 pim rp-address <ipv6-address> <acl_name>
```

If applied without any access list, the configured RP serves all the ff00::/8 space with the exception of only node/link local groups, SSM groups, and embedded RP address range.

The example configuration of three different RPs for different address ranges is shown below:

```
ipv6 pim rp-address 2001:efef:14:5145::145 range1
ipv6 pim rp-address 2001:cdcd:10A:6802::1
ipv6 pim rp-address 2001:abba:E000:501::2 range2
!
ipv6 access-list range1
 permit ipv6 any FF0B::/16
 permit ipv6 any FF1B::/16
 permit ipv6 any FF3B::/16
!
ipv6 access-list range2
 permit ipv6 any FF3E:30:2001:abba:1:FFFF::/96
```

Boot-Strap Router (BSR)

BSR is first specified in RFC2362; its capability is expanded in the new draft “draft-ietf-pim-sm-bsr” [\[17\]](#) to cover the possibility for administrative scoping in PIM domains using BSR.

Currently, Cisco IOS Software only forwards BSR messages (does not interpret them and does not learn the RP information from them). The BSR messages known to the router can be displayed using:

```
6net> show ipv6 pim bsr
```

PIMv2 BSR information

```
BSR Address: 2001:630:D0:130::1
Uptime: 3d07h, BSR Priority: 10, Hash mask length: 0
RPF: FE80::208:E2FF:FE3A:8200, POS3/0
Expires: 00:02:05
```

The full support for [\[17\]](#) is under development. Information on the BSR EFT (Early Field Trial) implementation is available in [Appendix B](#).

Embedded RP

Embedded RP is a new [\[3\]](#), IPv6-specific mechanism of learning the RP information (the group to RP mappings) using the encoding of the RP IPv6 address inside the multicast group destination address. The encoding has been described in detail in [Embedded RP Groups](#). This mechanism removes the need for complicated RP distribution procedures—every router on the multicast data path simply learns the RP automatically from the group address.

The implementation of embedded RP in the whole network moves the control over the selection of the RP IP address from the network administrator to the end user—the end user can dictate to the network what IP address should be the RP by just injecting multicast data with a conveniently formatted multicast group address. Draft [\[3\]](#) restricts the maximum prefix length embedded in the IP address to 64 bits. The RP address must always contain these maximum nonzero 64 bits, then the next minimum of 60 bits must always be zero and the last 4 bits are determined using the RPad field.

When designing a unicast IP addressing plan of an IPv6 Multicast-enabled network infrastructure, the network administrator must take the embedded RP addressing into account and make sure that physical interfaces of any router will never be eligible to function as the embedded RP. This can be achieved by making any of the 65th to 124th bits of the IP address nonzero. Such an addressing plan will protect the network infrastructure against

any kind of denial of service attacks (intended or accidental) by, for example, a high rate of PIM register packets forwarded to some low-speed physical interfaces, and will place the control over the RP location back into the hands of the network administrator by using appropriately selected loopback interfaces to serve as the RP in the embedded RP address ranges.

PIM Topology Information

The new PIM draft [\[14\]](#) defines a new concept as compared to RFC2362 regarding the PIM tree-building procedures and the corresponding state maintenance. It is reflected mainly in the definition of immediate and inherited outgoing interface lists (OILs) for the PIM states:

Immediate OIL—The list of interfaces added to a PIM topology table entry as a direct result of the receipt of explicit PIM Joins (from downstream routers) and MLD responses (for directly connected hosts).

Inherited OIL—The list of interfaces copied to an (S,G) PIM topology table entry from the Immediate OIL of the parent (*,G) PIM topology table entry (these interfaces are inherited from the [* ,G] Immediate OIL). In addition, there are cases where an (S,G,RPT) entry can “inherit” interfaces from a parent (*,G) Immediate OIL. The exact definition and usage of the Inherited OIL is beyond the scope of this document but is available in Section 4.1.6 of the PIM-SM Draft [\[14\]](#).

The combination of both Immediate and Inherited interfaces of a PIM topology table is used to build the multicast routing table (MRIB), which, in turn, is used to build the forwarding table (MFIB).

The above distinction has significant impact on OIL creation rules and also on the information provided by Cisco IOS Software commands, especially as compared with the IPv4 alternatives.

Immediate OIL

The immediate OIL can be checked using the following show commands:

```
show ipv6 mroute
show ipv6 pim topology
```

Note: The mroute command is present only for the compatibility with IPv4 outputs as the actual “mroute” data structure has been replaced by the combination of the topology table, the MRIB, and the MFIB data structures.

The following outputs have been taken from the Cisco 7200-2 router on [Figure 3](#) simulating a last-hop router with directly connected receivers for the FF05:1::5 group.

The relevant configuration on router 7200-2 is as follows:

```
interface Loopback1
  ipv6 address 2001:abcd:14::6/128
  ipv6 mld join-group FF05:1::5
!
interface GigabitEthernet0/0
  ipv6 address 2001:abcd:0:4::2/64
  ipv6 mld static-group FF05:1::5
```

```
7200-2>show ipv6 mroute ff05:1::5
```

Multicast Routing Table

Flags: D - Dense, S - Sparse, B - Bidir Group, s - SSM Group,
C - Connected, L - Local, I - Received Source Specific Host Report,
P - Pruned, R - RP-bit set, F - Register flag, T - SPT-bit set,
J - Join SPT
Timers: Uptime/Expires
Interface state: Interface, State

(*, FF05:1::5), 00:18:30/never, RP 2001:abcd:14::2, flags: SCLJ
Incoming interface: Tunnel10
RPF nbr: FE80::C0A8:F01
Outgoing interface list:
GigabitEthernet0/0, Forward, 00:18:30/never
Loopback1, Forward, 00:18:30/never

(2001:abcd:0:D::64, FF05:1::5), 00:16:02/00:01:28, flags: SJT
Incoming interface: Tunnel10
RPF nbr: FE80::C0A8:F01
Outgoing interface list: Null

The new command with extended information is:

7200-2>show ipv6 pim topology ff05:1::5

IP PIM Multicast Topology Table

Entry state: (*S,G)[RPT/SPT] Protocol Uptime Info

Entry flags: KAT - Keep Alive Timer, AA - Assume Alive, PA - Probe Alive,
RA - Really Alive, LH - Last Hop, DSS - Don't Signal Sources,
RR - Register Received, SR - Sending Registers, E - MSDP External,
DCC - Don't Check Connected

Interface state: Name, Uptime, Fwd, Info

Interface flags: LI - Local Interest, LD - Local Disinterest,
II - Internal Interest, ID - Internal Disinterest,
LH - Last Hop, AS - Assert, AB - Admin Boundary

(*,FF05:1::5)
SM UP: 00:20:02 JP: Join(00:00:16) Flags: LH
RP: 2001:abcd:14::2
RPF: Tunnel10,FE80::C0A8:F01
GigabitEthernet0/0 00:20:02 fwd LI LH
Loopback1 00:20:02 fwd LI II LH

(2001:abcd:0:D::64,FF05:1::5)

SM SPT UP: 00:17:34 JP: Join(00:00:16) Flags: KAT(00:03:26) RA

RPF: Tunnel10,FE80::C0A8:F01

No interfaces in immediate olist

SM—Indicates sparse mode protocol, in which the group operates

SM SPT—Indicates the shortest path tree created by SM

SM RPT—Indicates the (S,G,RPT) state

LH—Last Hop—Appears on interfaces where local MLD join has been received

LI—Local Interest—Appears on interfaces with configured static MLD join

II—Internal Interest—Appears on interfaces with configured MLD join-group, meaning that the router itself is receiving the data

The two outputs already demonstrate the significance of the immediate OILs—under the old RFC2362 (and common in IPv4 Multicast) rules, the OIL of the (*,G) state should have been copied into the (S,G) state. While the new draft [\[14\]](#) specifies that the immediate OIL contain only interfaces where explicit PIM or MLD joins have been received. The SPT tree and the corresponding (S,G) state have been created, but since no source-specific join has been received, the OIL of the (S,G) state remains empty.

Inherited OIL

The full list of interfaces, which will be used for forwarding the multicast data, can be seen using the show ipv6 mrib route command (the output taken from 7200-2 router at the same time as the outputs in [Immediate OIL](#)).

7200-2>show ipv6 mrib route ff05:1::5

IP Multicast Routing Information Base

Entry flags: L - Domain-Local Source, E - External Source to the Domain,

C - Directly-Connected Check, S - Signal, IA - Inherit Accept, D - Drop

Interface flags: F - Forward, A - Accept, IC - Internal Copy,

NS - Negate Signal, DP - Don't Preserve, SP - Signal Present,

II - Internal Interest, ID - Internal Disinterest, LI - Local Interest,

LD - Local Disinterest

(*,FF05:1::5) RPF nbr: FE80::C0A8:F01 Flags: C

GigabitEthernet0/0 Flags: F NS LI

Tunnel10 Flags: A NS

Loopback1 Flags: F IC NS II LI

(2001:abcd:0:D::64,FF05:1::5) RPF nbr: FE80::C0A8:F01 Flags:

Tunnel10 Flags: A

GigabitEthernet0/0 Flags: F NS

Loopback1 Flags: F IC NS

A—Accept—Indicates an interface where the data will be accepted. This is normally the RPF interface toward the RP in the case of the (*,G) entry, or the source in the case of the (S,G) entry. However, this is not a rule, as there are other conditions where data will be accepted on a non-RPF interface such as in the case of Bidir PIM.

F—Forward—Indicates interfaces where the data will be forwarded (local interest or PIM joins received).

IC—Internal Copy—Packet will be copied to the router itself on an interface with configured MLD join-group

In the case of the **show ipv6 mrib route** command, both (*,G) and (S,G) states contain all the interfaces (both immediate and inherited) where the data will be forwarded.

MULTICAST DATA FORWARDING

IPv6 Multicast builds a forwarding table similar to the unicast forwarding table used by Cisco Express Forwarding. However, unlike the unicast forwarding table, the multicast forwarding table requires three steps to be built:

1. Calculate unicast routing table
2. Create PIM tree/topology tables
3. Calculate forwarding tables

The content of the Multicast Forwarding Table (MFIB) can be displayed on all platforms using:

```
show ipv6 mfib
```

The multicast forwarding can be disabled (only on the software based platforms) using configuration option:

```
no ipv6 mfib
```

By default, forwarding is enabled when multicast routing is configured and does not appear in the configuration. When explicitly disabled, the router stops data forwarding but still fully runs PIM.

The status of MFIB can be monitored using:

```
7500>show ipv6 mfib interface
```

IPv6 Multicast Forwarding (MFIB) status:

Configuration Status: enabled, mode: distributed

Operational Status: running in distributed mode

MFIB interface	status	CEF-based output
		[configured, available]
FastEthernet4/0/0	up	[yes, ?]
FastEthernet4/0/1	up	[yes, ?]
POS4/1/0	up	[yes, ?]
Loopback1	up	[yes, ?]
Loopback100	up	[yes, ?]

```
Tunnel0      up      [yes      ,?      ]
Tunnel1      up      [yes      ,?      ]
Tunnel2      up      [yes      ,?      ]
```

The output and provided information differs on different platforms, depending on if forwarding happens in software/hardware and if it is centralized/distributed.

Software-Based Platforms

The Cisco 72xx and 75xx platforms are examples of Cisco routers that use centralized software switching for IPv6 Multicast. However, Cisco 7500 routers containing Versatile Interface Processor (VIP Version 2 and higher) interface cards can be configured for distributed forwarding using the following commands (IPv4 distributed Cisco Express Forwarding is required to be enabled as well):

```
ip cef distributed
!
ipv6 unicast-routing
ipv6 cef distributed
ipv6 multicast-routing
```

If running in distributed mode, the RP forwarding and LC forwarding information ([Hardware-Based Platforms](#)) appears in the outputs below.

Full MFIB

The keyword “verbose” adds only the “process switching” information here:

```
7200-2>show ipv6 mfib ff05:1::5 verbose
```

IP Multicast Forwarding Information Base

Entry Flags: C - Directly Connected, S - Signal, IA - Inherit A flag,
AR - Activity Required, D - Drop

Forwarding Counts: Pkt Count/Pkts per second/Avg Pkt Size/Kbits per second

Other counts: Total/RPF failed/Other drops

Interface Flags: A - Accept, F - Forward, NS - Negate Signalling

IC - Internal Copy, NP - Not platform switched

SP - Signal Present

Interface Counts: FS Pkt Count/PS Pkt Count

(* ,FF05:1::5) Flags: C

Forwarding: 0/0/0/0, Other: 4246/4246/0

Tunnel10 Flags: A NS

GigabitEthernet0/0 Flags: F NS

Pkts: 0/0 (process switching)

Loopback1 Flags: F IC NS

Pkts: 0/0 (process switching)

(2001:abcd:0:D::64,FF05:1::5) Flags:

```
Forwarding: 141504/100/186/145, Other: 0/0/0
Tunnel10 Flags: A
GigabitEthernet0/0 Flags: F NS
  Pkts: 0/94682 (process switching)
Loopback1 Flags: F IC NS
  Pkts: 0/141504 (process switching)
```

Under normal PIM-SM (when not using Bidir PIM), flag A indicates the RPF interface for the corresponding *,G or S,G entry; flag F on the interface indicates an outgoing forwarding interface. Flag IC (Internal Copy) indicates interfaces with MLD join-group command (the router itself receiving data not only forwarding them).

Currently Active Sources

```
7200-2>show ipv6 mfib ff05:1::5 active
```

Active IPv6 Multicast Sources - sending >= 4 kbps

Group: FF05:1::5

Source: 2001:abcd:0:D::64

Rate: 100 pps/145 kbps(1sec), 142 kbps(last 1537 sec)

Packet Counts

```
7200-2>show ipv6 mfib ff05:1::5 count
```

IP Multicast Statistics

62 routes, 14 groups, 0.07 average sources per group

Forwarding Counts: Pkt Count/Pkts per second/Avg Pkt Size/Kilobits per second

Other counts: Total/RPF failed/Other drops(OIF-null, rate-limit etc)

Group: FF05:1::5

RP-tree: Forwarding: 0/0/0/0, Other: 4246/4246/0

Source: 2001:abcd:0:D::64, Forwarding: 149117/100/186/145, Other: 0/0/0

Tot. shown: Source count: 1, pkt count: 149117

Hardware-Based Platforms

The outputs below cover the GSR hardware-based distributed switching platform. These outputs differ from the software-based platforms in the following information:

RP Forwarding—Displays traffic forwarded by the central GSR Route Processor (GRP). This traffic is rate limited to 1000 packets per second (pps) in total (for all traffic forwarded by the whole GSR). Forwarding on the GRP happens always when tunnels are configured on the router and appear in the outgoing interface lists.

LC Forwarding—Displays traffic forwarded by the particular line-card CPU or a VIP card, in the case of distributed switching on Cisco 7500 routers.

HW Forwd—Displays traffic forwarded in hardware on the line card.

Full MFIB

```
12000-1>show ipv6 mfib ff05:1::5
```

IP Multicast Forwarding Information Base

Entry Flags: C - Directly Connected, S - Signal, IA - Inherit A flag,
AR - Activity Required, D - Drop

Forwarding Counts: Pkt Count/Pkts per second/Avg Pkt Size/Kbits per second

Other counts: Total/RPF failed/Other drops

Interface Flags: A - Accept, F - Forward, NS - Negate Signalling

IC - Internal Copy, NP - Not platform switched

SP - Signal Present

Interface Counts: Distributed FS Pkt Count/FS Pkt Count/PS Pkt Count

(* ,FF05:1::5) Flags: C

RP Forwarding: 0/0/0/0, Other: 0/0/0

LC Forwarding: 0/0/0/0, Other: 0/0/0

POS1/0 Flags: A

Tunnell1 Flags: F NS

Pkts: 0/0/0

POS1/11 Flags: F NS

Pkts: 0/0/0

(2001:abcd:0:D::64,FF05:1::5) Flags:

RP Forwarding: 7816536/29/186/42, Other: 0/0/0

LC Forwarding: 0/0/0/0, Other: 0/0/0

HW Forwd: 0/0/0/0, Other: 744/744/0

POS1/0 Flags: A F

Pkts: 0/0/0

Tunnell1 Flags: F NS

Pkts: 0/0/7816612

Currently Active Sources

```
12000-1>show ipv6 mfib ff05:1::5 active
```

Active IPv6 Multicast Sources - sending >= 4 kbps

Group: FF05:1::5

Source: 2001:abcd:0:D::64

RP Rate: 30 pps/43 kbps(1sec)

LC Rate: 0 pps/0 kbps(1sec)

HW Rate: 0 pps/0 kbps(1sec)

Packet Counts

```
12000-1>show ipv6 mfib ff05:1::5 count
```

IP Multicast Statistics

59 routes, 11 groups, 0.09 average sources per group

Forwarding Counts: Pkt Count/Pkts per second/Avg Pkt Size/Kilobits per second

Other counts: Total/RPF failed/Other drops(OIF-null, rate-limit etc)

Group: FF05:1::5

RP-tree:

RP Forwarding: 0/0/0/0, Other: 0/0/0

LC Forwarding: 0/0/0/0, Other: 0/0/0

Source: 2001:abcd:0:D::64,

RP Forwarding: 7838037/30/186/43, Other: 0/0/0

LC Forwarding: 0/0/0/0, Other: 0/0/0

HW Forwd: 0/0/0/0, Other: 744/744/0

Tot. shown: Source count: 1, pkt count: 7838037

MFIB Drop Entries

The outputs of MFIB entries also contain “drop” entries, which are present in order to drop traffic in Cisco Express Forwarding or in hardware.

All the traffic matching these entries will be dropped. These entries are identified using:

Drop Flag at the entry level:

```
12000-1>show ipv6 mfib
```

```
(* ,FF00::/15) Flags: D
```

RP Forwarding: 0/0/0/0, Other: 0/0/0

LC Forwarding: 0/0/0/0, Other: 0/0/0

Only one of the A or F flags is present at the interface level in the output of the corresponding entry:

```
12000-1>show ipv6 mfib FF05:1::5
```

```
(2001:abcd:0:D::64,FF05:1::5) Flags:
```

RP Forwarding: 12376/30/186/44, Other: 0/0/0

LC Forwarding: 0/0/0/0, Other: 5/5/0

FastEthernet4/0/1 Flags: A

meaning that the RPF interface is known, but no outgoing interface exists.

```
12000-1>show ipv6 mfib FF05:1::5
```

```
(2001:abcd:0:D::64,FF05:1::5) Flags:
```

```
RP Forwarding: 12376/30/186/44, Other: 0/0/0
LC Forwarding: 0/0/0/0, Other: 5/5/0
Loopback1 Flags: F IC NS
Pkts: 12375/0/0
Tunnel0 Flags: F NS
Pkts: 0/0/12376
```

meaning that outgoing interfaces are available (received joins) but no RPF interface exists (due to some routing failure).

ADMINISTRATIVE DATA SCOPING

Administrative Boundary

The administrative boundaries based on the configured group ranges (using access lists) and applied to the interfaces in the IPv4 style are under development for IPv6.

IPv6 Multicast Scopes

The deployment and implementation of IPv6 Multicast address scopes as defined in [Generic Multicast Group Addresses Definition](#) are under discussion at the IETF. The situation is more complicated by the unicast scoping (link local, site local) of the unicast source addresses, which needs to be taken into account as well—multicast data with a globally scoped group address while with site local unicast source address must not be forwarded out of the site scope. The unicast site local addressing is also under discussion at the IETF and might be removed.

A limited implementation of scopes/zones is available in Cisco IOS Software, but at the moment it is still hidden and has not been tested yet.

USER TO NETWORK SIGNALING

The IPv6 alternative of Internet Group Management Protocol (IGMP) is Multicast Listener Discovery (MLD) protocol specified as part of the IPv6 Internet Control Message Protocol (ICMP)—RFC2463. The protocol packets are sent using the link local addresses as source addresses and the destination address is typically the multicast group address which the message concerns to—no MLD-specific group address is allocated. The general queries need to be sent to the “all nodes” destination group address—FF02::1. The done message (IGMP Leave message) is sent to all routers address—FF02::2.

MLD packets are sent with the “Router Alert” option in the IPv6 header. Based on this option, the router interprets the packets as control packets and does not create any PIM states for it. If some (S,G) states appear having link local IPv6 addresses as the source addresses, it is an indication of a bug in the host software (missing Router Alert option), which needs to be patched (appears in Solaris 2.7, for example).

MLDv1

Multicast Listener Discovery Version 1 is specified by RFC2710. It provides capabilities equivalent to IGMPv2—RFC2236.

MLDv2

MLDv2 [\[18\]](#) extends the v1 capabilities the similar way as IGMPv3 (RFC3376)—it adds the possibility to request data only from certain sources, and groups not only the data for the whole group as v1 does. This is necessary for full support of SSM.

MLD Compatibility

MLDv2 router in the presence of the v1 host has to operate in the v1 compatibility mode. The compatibility status is kept based on the multicast address (multicast group addresses that have only the v2 host joined will operate in v2 mode). The General Query is always sent as v2, regardless of the compatibility mode.

Configuring and Monitoring MLD

Cisco IOS Software implements MLDv2:

Enabling MLD

MLD is enabled by default when multicast routing is enabled. The operation of MLD is not dependent on PIM as it was in the older Cisco IOS Software IPv4 Multicast code. It can be enabled separately using:

```
ipv6 mld router
```

The command appears in the configuration only in its negative form. An explicit **no ipv6 pim** command does not disable MLD, but disables multicast data forwarding over that interface (MLD messages are still processed):

The relevant configuration on the 7200-2 router (see [Figure 3](#)):

```
interface GigabitEthernet0/0
  ipv6 address 2001:abcd:0:4::2/64
  no ipv6 pim
```

```
7200-2>show ipv6 mld interface
```

```
GigabitEthernet0/0 is up, line protocol is up
  Internet address is FE80::208:A4FF:FEA7:A808/10
  MLD is enabled on interface
  Current MLD version is 2
  MLD querying router is FE80::208:A4FF:FEA7:A808 (this system)
```

Monitoring MLD

The MLD operation can be checked using **show ipv6 mld interface** command:

```
7200-2>show ipv6 mld interface
```

```
7200-2>GigabitEthernet0/0 is up, line protocol is up
  Internet address is FE80::208:A4FF:FEA7:A808/10
  MLD is enabled on interface
  Current MLD version is 2
  MLD query interval is 125 seconds
  MLD querier timeout is 255 seconds
  MLD max query response time is 10 seconds
  Last member query response interval is 1 seconds
```

Inbound MLD access group is:
MLD activity: 6 joins, 2 leaves
MLD querying router is FE80::208:A4FF:FEA7:A808 (this system)

Restricting Access

MLD access lists can be configured to allow access only to certain groups or certain sources (and groups) using:

```
ipv6 mld access-group <acl_name>
```

interface level command.

Deny all sources for a group G:

```
deny any host G  
permit any any
```

Deny particular source and group:

```
deny host S host G  
permit any host G
```

ENABLING AND DISABLING MULTICAST COMPONENTS

This chapter summarizes the available Cisco IOS Software commands regarding enabling and disabling various multicast components in the global configuration mode and interface modes, and describes the dependencies between the various component capabilities and the Cisco IOS Software commands.

Enabling IPv6 Multicast Routing

IPv6 Multicast routing is enabled in the global configuration mode by the command:

```
ipv6 multicast routing
```

This command enables:

- PIM processing on all interfaces
- Multicast data forwarding on all interfaces (MFIB)
- MLD processing on all interfaces

On software-based platforms, multicast processing will be by default enabled by this command on the central processor. If the platform supports distributed processing (RSP/VIP) it must be explicitly configured, according to [Software-Based Platforms](#).

Hardware-based platforms will be enabled by default for hardware processing with the exceptions mentioned in [Hardware-Based Platforms](#).

PIM

PIM-enabled interface means that PIM messages will be sent and processed on the interface. By default, all IPv6-enabled interfaces become PIM-enabled when multicast is enabled globally.

IPv6 PIM can be disabled on a per-interface basis using:

```
no ipv6 pim
```

This command disables:

- PIM processing on the interface
- Data forwarding on the interface (MFIB)

MLD stays operational on the interface, but the received MLD messages do not cause the interface to appear in the OIL and data does not get forwarded to it.

The status of interfaces can be checked using:

```
7500>show ipv6 pim interface
```

```
6net-7500>sh ipv6 pim int
```

Interface	PIM	Nbr	Hello	DR
		Count	Intvl	Prior

FastEthernet4/0/1	on	1	30	1
-------------------	----	---	----	---

Address: FE80::207:ECFF:FEC2:6881

DR : this system

POS4/1/0	on	1	30	1
----------	----	---	----	---

Address: FE80::207:ECFF:FEC2:6820

DR : FE80::208:E2FF:FE3C:300

ATM5/0/0	off	0	30	1
----------	-----	---	----	---

Address: ::

DR : not elected

Loopback1	on	0	30	1
-----------	----	---	----	---

Address: FE80::207:ECFF:FEC2:6820

DR : this system

Tunnel0	off	0	30	1
---------	-----	---	----	---

Address: FE80::207:ECFF:FEC2:6820

DR : not elected

MFIB

Global Configuration Mode

MFIB can be disabled globally using in global configuration mode:

```
no ipv6 mfib
```

This command disables any multicast data forwarding on the device but allows PIM and MLD processing on all interfaces.

Interface Mode

Multicast forwarding in Cisco Express Forwarding can be disabled on a per-interface basis using the interface level:

```
no ipv6 mfib-cef
```

This command disables multicast data forwarding on the interface in interrupt mode (causing the traffic being process switched) while leaving PIM and MLD operational.

The multicast data forwarding status of interfaces can be checked using:

```
7500>show ipv6 mfib interface
```

IPv6 Multicast Forwarding (MFIB) status:

Configuration Status: enabled, mode: distributed

Operational Status: running in distributed mode

MFIB interface	status	CEF-based output
		[configured,available]
FastEthernet4/0/1	up	[yes ,?]
POS4/1/0	up	[yes ,?]
Loopback1	up	[yes ,?]
Tunnel0	up	[yes ,?]

The command currently does not check the status of MFIB on the interface (if MFIB gets disabled due to **no ipv6 pim** command, it does not get reflected in the output). The MFIB interface is down in this output only if the line protocol on the interface is down.

MLD

MLD is enabled on all interfaces when multicast routing is enabled globally. The router side of MLD (querying) can be disabled on a per-interface basis by the use of the following command:

```
no ipv6 mld router
```

Note: That if MLD is disabled in this manner, the interface will still process the incoming MLD host messages.

The status of interfaces can be checked using:

```
7500>show ipv6 mld interface
```

```
FastEthernet4/0/1 is up, line protocol is up
Internet address is FE80::207:ECFF:FEC2:6881/10
MLD is enabled on interface
Current MLD version is 2
MLD query interval is 125 seconds
MLD querier timeout is 255 seconds
MLD max query response time is 10 seconds
Last member query response interval is 1 seconds
Inbound MLD access group is:
MLD activity: 400 joins, 393 leaves
MLD querying router is FE80::207:ECFF:FEC2:6881 (this system)
```

```
Loopback1 is up, line protocol is up
  Internet address is FE80::207:ECFF:FEC2:6820/10
  MLD is enabled on interface
  Current MLD version is 2
  MLD query interval is 125 seconds
  MLD querier timeout is 255 seconds
  MLD max query response time is 10 seconds
  Last member query response interval is 1 seconds
  Inbound MLD access group is:
  MLD activity: 9 joins, 2 leaves
  MLD querying router is FE80::207:ECFF:FEC2:6820 (this system)
Tunnel0 is up, line protocol is up
  Internet address is FE80::207:ECFF:FEC2:6820/10
  MLD is disabled on interface
```

DEPLOYMENT SCENARIOS

Single-Domain Multicast Scenarios

Single-domain scenarios represent cases where the whole network is under the control of one administration.

Congruent Native Topology

In this simplest case, where the whole network is fully IPv6- and multicast-enabled, no special configuration is required. In this case, normal unicast routing will fully take care of any RPF check requirements:

1. Enable globally multicast routing
`ipv6 multicast-routing`

This enables PIM, MLD, and multicast data forwarding on all interfaces enabled for IPv6.

2. Decide for the RP distribution method

a) Configure the RP address(es) statically on all routers:

```
ipv6 pim rp-address 2001:efef:14:5145::145 range1
ipv6 pim rp-address 2001:cdcd:10A:6802::1
ipv6 pim rp-address 2001:abba:E000:501::2 range2
```

```
ipv6 access-list range1
  permit ipv6 any FF0B::/16
  permit ipv6 any FF1B::/16
  permit ipv6 any FF3B::/16
```

```
!
```

```
ipv6 access-list range2
  permit ipv6 any FF3E:30:2001:abba:1:FFFF::/96
```

b) Deploy only SSM—This does not require any configuration. SSM FF3x::/32 range of multicast group addresses must be in use only.

c) Deploy embedded RP—No configuration is necessary across the network. Embedded RP multicast group addresses must be in use. The RP itself needs to be configured as follows:

```
ipv6 pim rp-address 2001:efab:0:FE::1 RP-embedded
```

```
ipv6 access-list RP-embedded  
  permit ipv6 any FF7B:140:2001:efab:0:FE::/96
```

where 2001:efab:0:FE::1 is the local router address.

Noncongruent Native Topology

To create different forwarding paths for unicast and multicast data on the same physical infrastructure without encapsulating multicast data in unicast IPv6 packets (tunneling), the router needs to hold a separate unicast routing table that is not visible for forwarding unicast packets but can be only used for RPF checks. In IPv6 Multicast this is possible only using static routes ([Static Routing](#)) or MP-BGP ([BGP for IPv6](#) and [IPv6 MP-BGP CONFIGURATION EXAMPLE](#)). This limits its deployment only to small networks (static routing) or large BGP backbones.

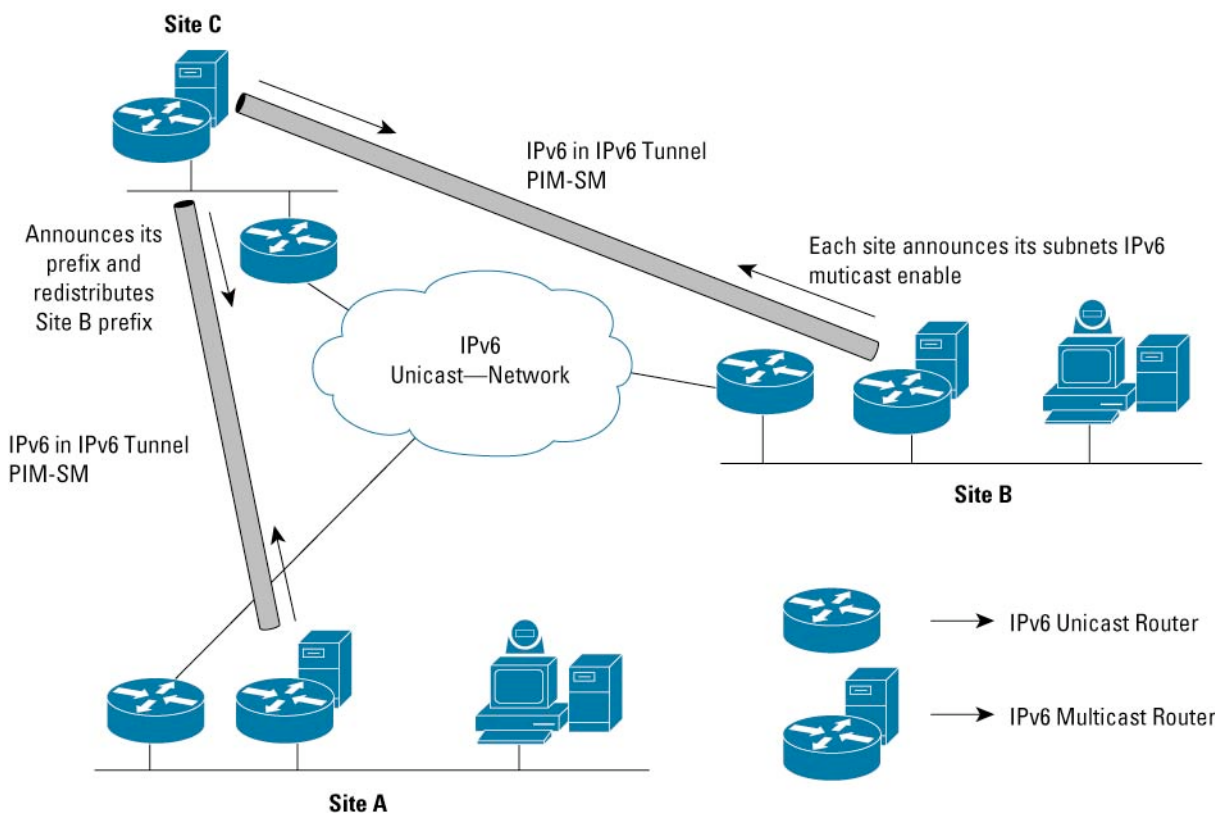
Noncongruent Tunneled Topology

Noncongruent topology means different multicast topology to the physical layout of the network for several possible reasons:

1. No support for IPv6 Multicast in some parts of the network
2. No support for IPv6 at all in some parts of the network

These cases do not differ much in configuration, and need to be solved with tunneling. The cloud on [Figure 4](#) can represent either an IPv6 or IPv4 network.

Figure 4. Noncongruent Topology Scenario (Picture used with the permission of the authors of [19])



Selecting Tunnel Endpoints

The design differs considerably depending on where it is possible to terminate the tunnels:

1. Separate multicast routers are available (Figure 4)—this allows creation of a unicast routing table for RPF checks using a dynamic routing protocol. Care must be taken that the protocols running on unicast-only routers and multicast-only routers do not interfere and do not exchange any routing information. The multicast router must only have routing information to the tunnel endpoints (/128 IPv6 or /32 IPv4 prefixes) or a static default route. The tunnel endpoints must never be unicast-routed over the tunnel interface itself, but through some physical interface.
2. The access router on each site has to run both unicast and multicast—no possibility to put the two functions on two physical boxes. In that case, the only feasible possibility to forward unicast over the physical link and multicast over the tunnels between the sites is static routing (another option would be MP-BGP, but it is not typically run to the endsites) for RPF checks. The tunnel endpoints must never be resolved by unicast routing to use the tunnel interface itself, but through some physical interface. Each site needs to be configured with static routes pointing to the tunnels as follows (or more easily just with a multicast default route pointing to the tunnel):

Site A:

```
ipv6 route <prefix_site_B> tunnel0 <next_hop_address> multicast
ipv6 route <prefix_site_C> tunnel0 <next_hop_address> multicast

ipv6 route 2001::/16 Tunnel10 2002::1 multicast
```

Configuring Tunnels

1. IPv4 connectivity available:

```
!  
interface Tunnel10  
  description tunnel to 12000-1  
  no ip address  
  ipv6 address 2002::2/64  
  ipv6 enable  
  tunnel source 192.168.15.2  
  tunnel destination 192.168.15.1  
  tunnel mode ipv6ip  
!
```

If dynamic routing protocols can be run on the tunnel (separate multicast router case), the following commands (using RIP as an example) need to be added to the tunnel interface configuration and to the local site Ethernet:

Enable routing process (RIP) globally:

```
ipv6 router rip test
```

Enable interfaces for RIP

```
interface Tunnel10  
ipv6 rip test enable
```

```
interface Ethernet0/0  
ipv6 rip test enable
```

If the routing protocol allows, the Ethernet interface should be made passive, or a routing updates filter list should be applied to deny all incoming routes and avoid any interaction with the unicast routing on the other unicast-dedicated router. The tunnel endpoint 192.168.15.1 needs to be made reachable using the easiest IPv4 static routing.

2. IPv6 connectivity available

```
interface Tunnel111  
  no ip address  
  ipv6 address 2001:aaaa::2/64  
  tunnel source 2001:abcd:C::1  
  tunnel destination 2001:abcd:C::2  
  tunnel mode ipv6
```

Only the tunnel mode (encapsulation) and the source and destination fields need to be changed; the routing protocols configuration stays the same as in the previous case.

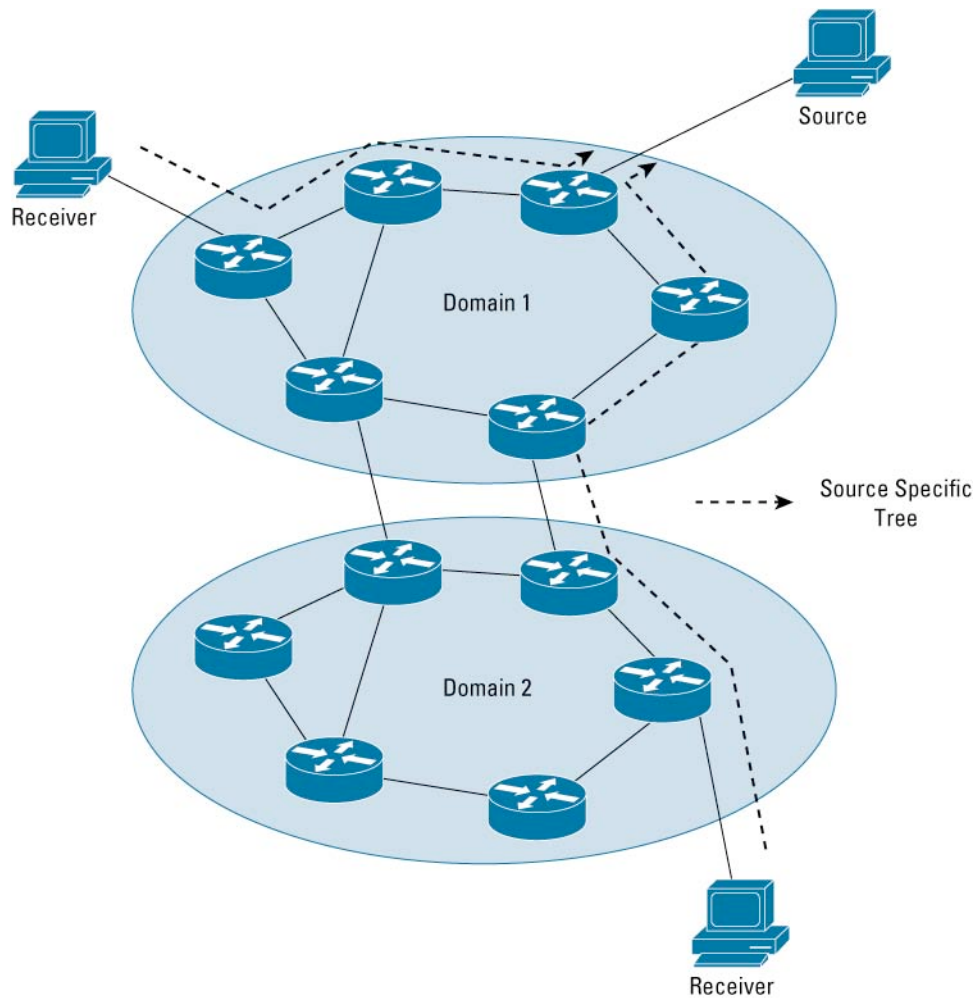
RP Distribution Method

The same methods for the RP distribution are available as discussed in [Congruent Native Topology](#).

SSM

SSM ([\[20\]](#), RFC3569) represents a subset of PIM-SM (RFC2362) that removes all procedures related to the RP and the Shared Tree and employs Shortest-Path Trees only using the source- and group-specific (S,G) procedures of PIM-SM. This means that the host application is required to know the source and group address of the content stream it wishes to receive and must signal this using IGMPv3 (IPv4) or MLDv2 (IPv6). (**Note:** The RP and the Shared Tree in PIM-SM perform network-based source discovery, freeing the host application from this task. However, this comes at the expense of increased complexity and is often unnecessary as many applications use a one-to-many multicast model where the application has immediate access to source information.) [Figure 5](#) is an example of SSM being used for interdomain multicast. Notice that only Shortest-Path Trees are being employed to deliver the content from the source.

Figure 5. SSM Example of Operation in the Interdomain Scenario



On the user side, SSM requires implementation of MLDv2. The network side is required to recognize that it should apply SSM procedures and originate only (S,G) PIM join directly to the source. In the IPv4 case, this recognition is based on the configured SSM address range with [\[21\]](#) as default. IPv6 strictly assigned in its addressing architecture discussed in [MULTICAST GROUP ADDRESSING](#) an address range for SSM only, and network devices must not apply any RP-specific procedures for these addresses.

The deployment of SSM is a fully valid option in both intra- and interdomain scenarios. The deployment scenario with SSM is quite simple—it only requires PIM enabled in the networks and the last-hop routers being aware of the SSM address ranges. The drawback is the potential growth of (S,G) forwarding states if the multicast deployment reaches the level of unicast. The deployment of SSM is still complicated with the lack of support for MLDv2 in the end-user IP stacks.

Interdomain Multicast Scenarios

Interdomain scenarios represent the case where the network is divided into several administrative domains, based only on the equipment ownership and not related to any protocol functions.

Interdomain multicast in IPv4 allowed each PIM domain to administer its own RP. In order to facilitate interdomain multicast, the Multicast Source Discovery Protocol (MSDP) was developed to exchange information about sources between PIM domains. This allowed each RP to issue PIM (S,G) joins to sources in other domains when having receivers for group G. The operation of the protocol has proven to be troublesome and not scalable for many applications; IETF has so far strictly refused to discuss this protocol for IPv6.

When speaking about interdomain in IPv6, the administrative domains are forced to share the RP unless other mechanisms are used. In the Internet environment, this administrative domain is typically also a BGP Autonomous System, and BGP will be (in most cases) the routing protocol run between the domains.

The IETF discussion about interdomain multicast in IPv6 now concentrates around SSM ([\[20\]](#), RFC3569), which could replace Any Source Multicast (ASM in the sense of the full RFC2362 implementation) using multiple one-to-many trees in order to create full mesh of any-to-any communications.

MP-BGP Deployment

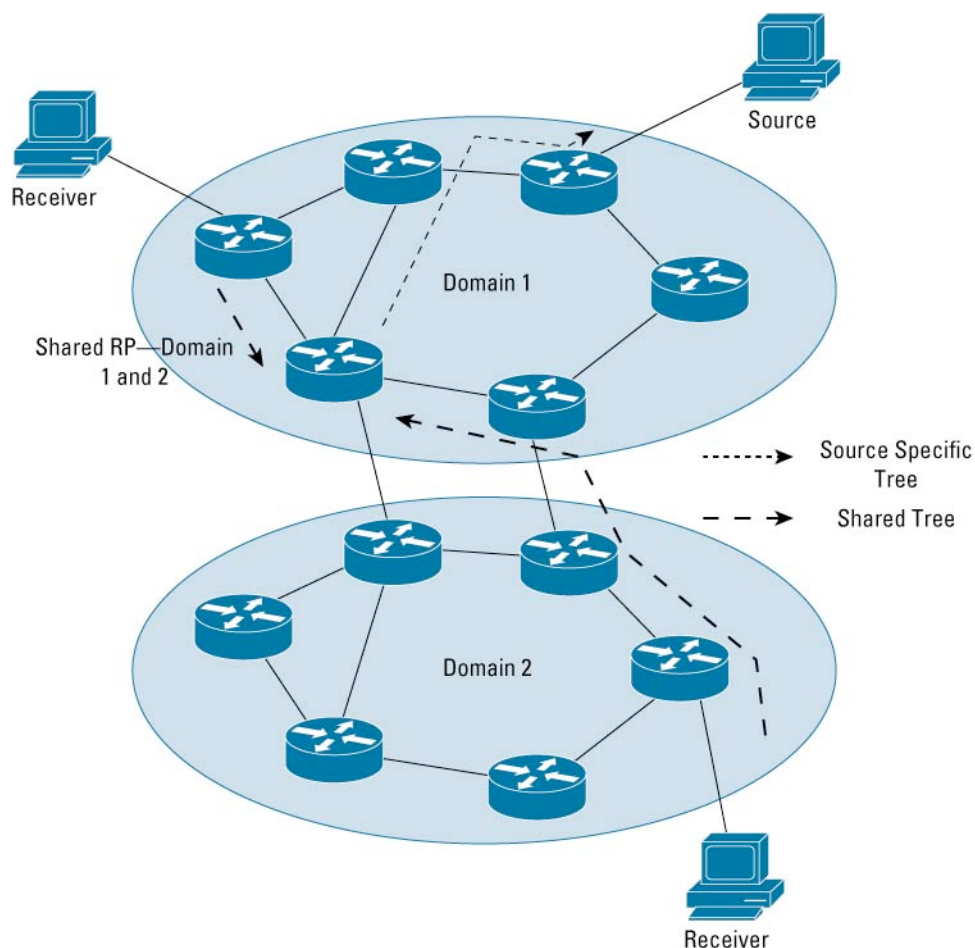
MP-BGP is not “a MUST” when the network administrator decides to run multicast. If the topology used to deliver unicast and multicast is congruent (or exactly the same), there is no need to enable MP-BGP. Congruent in terms of BGP means:

1. All BGP next hops are equal for all potential unicast and multicast BGP routes
2. The unicast and multicast BGP tables contain exactly the same prefixes and the same BGP attributes

Any Source Multicast

[Figure 6](#) is an example of Any Source Multicast (ASM, RFC2362, [\[14\]](#)) operation spanning several administrative domains. In the IPv6 case, ASM requires all PIM domains to share one RP for a particular multicast group range, which has many technical and administrative problems. At present, there is a lack of experience and best practices in this area.

Figure 6. ASM Example of Operation in the Interdomain Scenario



Embedded Rendezvous Point

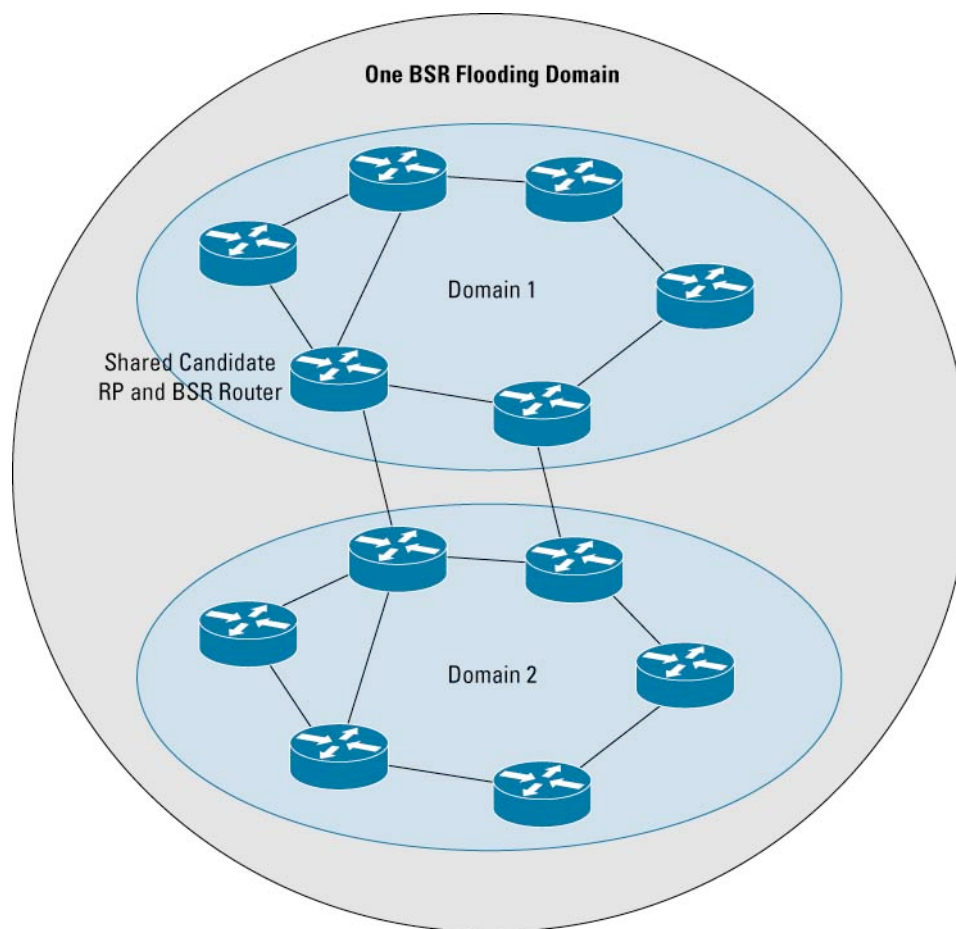
The procedures of [3] can be used in both intra- and interdomain scenarios, but the main idea is to resolve the necessity of the RP group to mapping distribution between domains. As opposed to SSM, all the network devices all the way must implement embedded RP.

Shared Static RP

In this scenario, all PIM domains of the internetwork need to be configured with at least one common RP to handle groups, which will have sources and receivers in all domains. Each domain can have its own set of RPs to handle group ranges specific for each of the local domains. The configuration of each of the network devices would be something like in the example in [Distributing RP Information](#). Potentially, the static configuration can be used for the interdomain-shared RP; in intradomain, it can be substituted by a dynamic distribution of the RP information (BSR, embedded).

The static configuration represents a certain administration burden, but if IP addressing is planned sensibly, the shared RP can be moved to different physical locations inside of the particular hosting domain without actually changing the IP address. The shared RP should be physically positioned in the backbone (if the whole internetwork has some hierarchy) so all the users of it have approximately the same physical distance to it. Static configuration also has the advantage of avoiding debugging of pure PIM control information distribution on top of multicast data distribution.

Figure 7. Shared BSR Flooding Domain Example

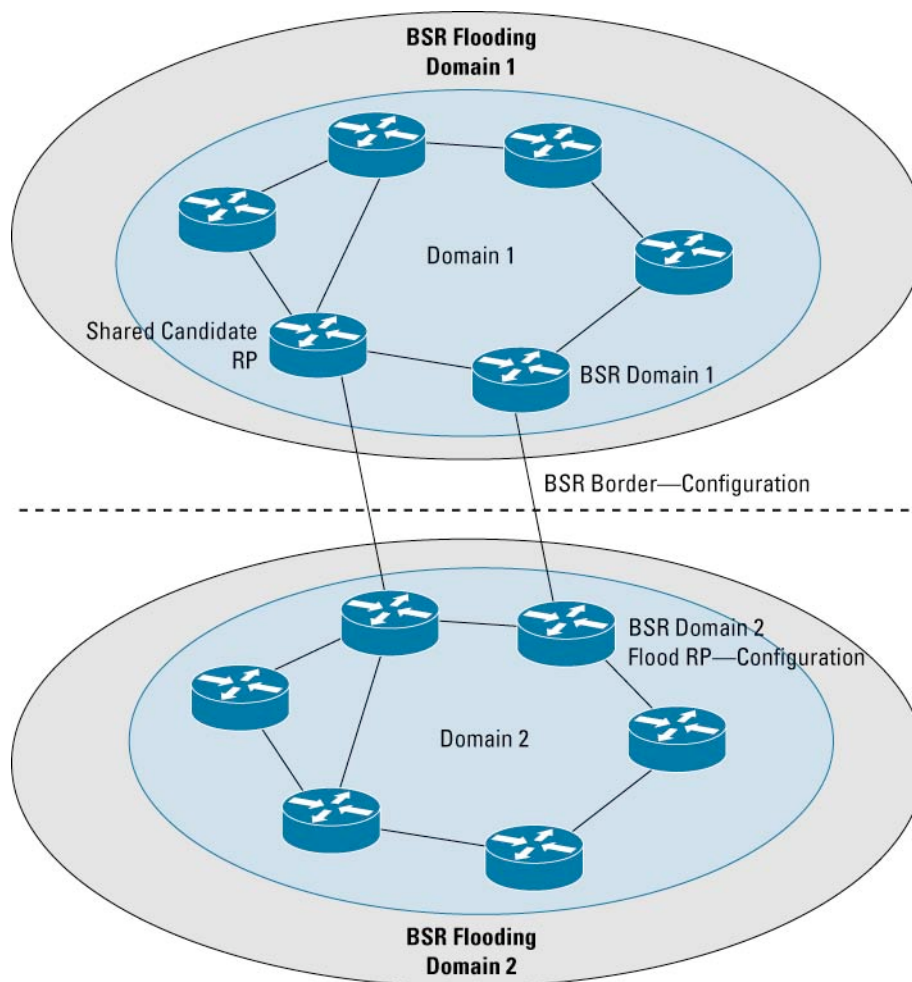


Shared BSR Domain

[Figure 7](#) depicts an example of a shared BSR domain, where RP information is flooded across multiple domains. In this scenario, the whole internetwork running multicast represents just one BSR domain—the identical BSR information is flooded freely without any constraints.

All the domains must share the exactly same set of RPs for all groups. The distribution of locally significant RP information would not be possible using BSR but only via static configuration. The implementation of the new BSR draft [\[17\]](#) might bring the possibility to use administrative scoping for BSR messages in this scenario to help to solve the problem above. This scenario represents a highly cooperative network, where one entity manages the operation of the whole internetwork with BSR routers and candidate RPs most probably placed in the interconnecting backbone. The distribution of the shared RP information is administratively flexible, but the debugging is more complicated—RP distribution and data distribution need to be debugged.

Figure 8. Constrained BSR Flooding Domains Example



Constrained BSR Domains

[Figure 8](#) shows an example of a “constrained” BSR domain. In this scenario, each PIM domain would also be a separate BSR domain. Each domain entirely blocks any BSR flooding/traffic to any other domain. The domains need to agree on a common set of RPs and group ranges, which will be flooded identically in each domain. The BSR routers need to allow a static configuration of RPs, which need to be advertised (without any BSR messages from the candidate RPs to the BSR routers). On top of that, each domain can flood its own information about local RPs inside of the domain.

This scenario also requires a certain level of administrative coordination, but decreases the size of the BSR domains and allows constrained debugging of BSR flooding to each of the domains with each domain responsible for the distribution inside of its infrastructure. It also gives each domain full flexibility to advertise locally significant RPs internally via BSR. (**Note:** Support of this scenario requires implementation of DDTS CSCeb73511 in Cisco IOS Software.)

IPv6 MP-BGP CONFIGURATION EXAMPLE

The BGP configuration of a core part of the network is practically identical to the unicast or IPv4 case. Some new BGP features (like peer-group templates or dynamic update groups) might not be available for the IPv6 address families and need to be checked on a case-by-case basis.

Internal BGP Core

The iBGP core requires the usual full mesh of peerings or any of the route-reflection or confederation scenarios common in general BGP.

Global BGP Process and Peer Definition

Each BGP speaker has to be configured explicitly with a 32 bit entity (sometimes called IPv4 address) in order to acquire the BGP router ID necessary in the BGP OPEN message (RFC1771) to establish the BGP session.

```
router bgp 65000
  bgp router-id 10.10.10.10
  neighbor INTERNAL_PEER peer-group
  neighbor INTERNAL_PEER password <removed>
  neighbor INTERNAL_PEER update-source Loopback0
  neighbor 2001:abcd:10::1 remote-as 65000
  neighbor 2001:abcd:10::1 peer-group INTERNAL_PEER
  neighbor 2001:abcd:12::1 remote-as 65000
  neighbor 2001:abcd:12::1 peer-group INTERNAL_PEER
  neighbor 2001:abcd:16::1 remote-as 65000
  neighbor 2001:abcd:16::1 peer-group INTERNAL_PEER
  neighbor 2001:abcd:17::1 remote-as 65000
  neighbor 2001:abcd:17::1 peer-group INTERNAL_PEER
  neighbor 2001:abcd:20::1 remote-as 65000
  neighbor 2001:abcd:20::1 peer-group INTERNAL_PEER
```

Activate Peers in Address Families

```
router bgp 65000
  address-family ipv6 multicast
    neighbor INTERNAL_PEER activate
    neighbor INTERNAL_PEER send-community
    neighbor INTERNAL_PEER advertisement-interval 1
    neighbor 2001:abcd:10::1 peer-group INTERNAL_PEER
    neighbor 2001:abcd:12::1 peer-group INTERNAL_PEER
    neighbor 2001:abcd:16::1 peer-group INTERNAL_PEER
    neighbor 2001:abcd:17::1 peer-group INTERNAL_PEER
    neighbor 2001:abcd:20::1 peer-group INTERNAL_PEER
  exit-address-family
!
address-family ipv6
  neighbor INTERNAL_PEER activate
```

```
neighbor INTERNAL_PEER send-community
neighbor INTERNAL_PEER advertisement-interval 1
neighbor 2001:abcd:10::1 peer-group INTERNAL_PEER
neighbor 2001:abcd:12::1 peer-group INTERNAL_PEER
neighbor 2001:abcd:16::1 peer-group INTERNAL_PEER
neighbor 2001:abcd:17::1 peer-group INTERNAL_PEER
neighbor 2001:abcd:20::1 peer-group INTERNAL_PEER
no synchronization
exit-address-family
```

External Peerings

Global Peers Definition

```
router bgp 65000
neighbor EXTERNAL_PEER peer-group
    neighbor 2001:cccc:0:B000::1 remote-as 65001
    neighbor 2001:cccc:0:B000::1 peer-group EXTERNAL_PEER
    neighbor 2001:cccc:0:B000::1 password <removed>
    neighbor 2001:abcd:14:200::2 remote-as 65002
    neighbor 2001:abcd:14:200::2 peer-group EXTERNAL_PEER
neighbor 2001:abcd:14:200::2 password <removed>
```

Activate Peers in Address Families

```
router bgp 65000
address-family ipv6 multicast
    neighbor EXTERNAL_PEER activate
    neighbor EXTERNAL_PEER send-community
    neighbor EXTERNAL_PEER advertisement-interval 5
    neighbor 2001:cccc:0:B000::1 peer-group EXTERNAL_PEER

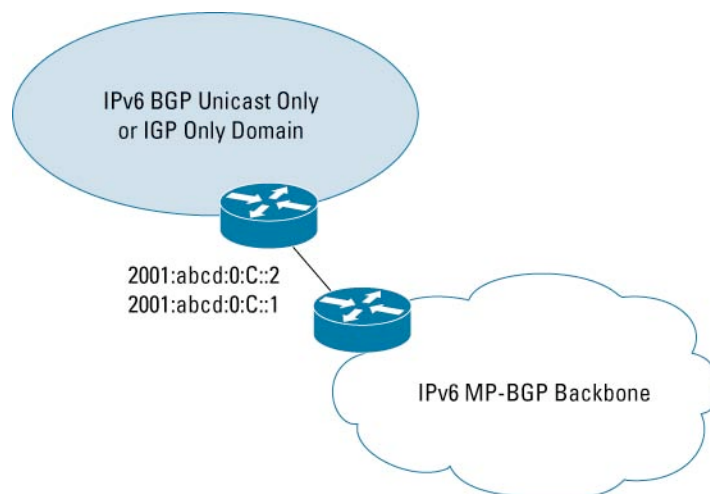
exit-address-family
!
address-family ipv6
    neighbor EXTERNAL_PEER activate
    neighbor EXTERNAL_PEER send-community
    neighbor EXTERNAL_PEER advertisement-interval 5
    neighbor EXTERNAL_PEER soft-reconfiguration inbound
    neighbor EXTERNAL_PEER route-map FROM-PEER in
    neighbor 2001:cccc:0:B000::1 peer-group EXTERNAL_PEER
    neighbor 2001:abcd:14:200::2 peer-group EXTERNAL_PEER
exit-address-family
```

Connecting Non-MP-BGP Domains

This example covers two cases of backbone clients where MP-BGP is not running between the two connected networks as is shown in [Figure 9](#).

1. The client runs only Internal Gateway Protocol (IGP) like IS-IS, RIP or OSPF
2. The client runs BGP but does not support MP-BGP

Figure 9. Non-MP-BGP Backbone Clients



In the first case the backbone border router needs to run the client's IGP on the link between the two and redistribute it into its MP-BGP on behalf of the client—see example in [Redistribute IGP](#).

In the second case, the client already originates his BGP prefixes, but only with SAFI=1 (unicast address family). The backbone border router needs to take these prefixes and not only propagate them with SAFI=1 but also inject them with SAFI=2 while keeping all other BGP attributes like AS-path. This case is covered by the example [Translate Updates from Non-MP-BGP Peer](#)—both keywords “unicast multicast” must be present in the translate-update command in order to propagate updates with both SAFI=1 and SAFI=2.

Redistribute IGP

The BGP neighbors specified in the multicast address family are internal MP-BGP peer towards the backbone routers. RIP process “test” (or any other user specified process name or routing protocol) needs to be enabled on the interface leading to the client and redistributed into the multicast address family.

```
address-family ipv6 multicast
  neighbor 2001:abcd:14::2 activate
  neighbor 2001:abcd:14::3 activate
  redistribute rip test
exit-address-family
```

Translate Updates from Non-MP-BGP Peer

In this case the client can run BGP, but only unicast address family. In the example below the border router has same internal BGP peers activated for both unicast and multicast address family as in the previous example. The 2001:abcd:0:C::2 peer is the unicast only and external peer and needs to be configured under the address family unicast with the translate-update command. To make this configuration effective and inject the updates with multicast address family the client BGP peer has to be activated under the multicast address family as well although it does not support it.

```
address-family ipv6
    neighbor 2001:abcd:0:C::2 activate
    neighbor 2001:abcd:0:C::2 translate-update ipv6 multicast unicast
    neighbor 2001:abcd:14::2 activate
    neighbor 2001:abcd:14::3 activate
exit-address-family
!
address-family ipv6 multicast
    neighbor 2001:abcd:0:C::2 activate
    neighbor 2001:abcd:14::2 activate
    neighbor 2001:abcd:14::3 activate
exit-address-family
```

APPENDIX A—PIM INTERFACE MODES AND GROUP MODES—V4 TO V6

PIM Group Modes

In Cisco IOS Software, each IP multicast group operates in exactly one mode of PIM. This group mode is determined by configuration and/or learned via dynamic protocols as described below. For IPv4 Multicast this is true starting with Cisco IOS Software version 12.0 and later—which introduced PIM version 2 for IPv4, prior Cisco IOS Software versions used exclusively PIM version 1. For IPv6 Multicast this is true for all versions of Cisco IOS Software supporting IPv6 Multicast (12.0(26)S, 12.2(18)S, 12.3(2)T and later).

Cisco IOS Software IPv4 Multicast supports PIM Sparse Mode, PIM SSM, PIM Dense Mode and Bidir PIM. Cisco IOS Software IPv6 Multicast supports PIM Sparse Mode, PIM SSM and Bidir PIM. The PIM mode of a group is determined as follows:

1. If a group falls into the SSM range, then PIM uses the SSM mode, which basically is PIM Sparse Mode but without any RP, any (*,G) state or (*,G) messages. A Cisco IOS Software router will ignore any RP configuration for groups in the SSM range.

In IPv4, the SSM range needs to be configured on a router via the **ip multicast ssm [default | range ...]** command. In IPv6 a group is implicitly recognized to be in the SSM range by the IPv6 Multicast group address formatting itself.
2. If a group is not SSM, but the router knows a RP for the group, then the mode of the group is determined by the mode of the RP (being either PIM-SM or Bidirectional PIM). A RP and its mode can be learned dynamically via the Bootstrap Router protocol of PIM (BSR), static configuration or AutoRP. BSR and static configuration are available for IPv4 and IPv6, AutoRP only for IPv4.
3. If Cisco IOS Software does not know a RP for a group, then it uses a default group mode. In IPv4, this default mode is PIM Dense Mode. In IPv6, this mode is PIM Sparse Mode—just that no RP is known. PIM Sparse Mode without a configured RP behaves very similar to PIM-SSM, the difference being is that there is (*,G) state, but it's RPF interface is Null—e.g.: dysfunctional, but as soon as the router learns an RP for the group again—for example via BSR—it will become fully operational again.

PIM Interface Modes

Cisco IOS Software IPv4 Multicast supports several PIM interface modes, and it has often been confusing for customers to understand how they relate to the PIM group mode. With IPv6 Multicast this distinction is gone, but how does that now compare to IPv4 Multicast? The following table lists all cases:

Group Mode	Cisco IOS Software IPv4 Multicast Interface Mode			Cisco IOS Software IPv6
	dense-mode	sparse-mode	sparse-dense-mode	
SSM group	Not supported	Works	Works	Works
Sparse group	Converted to Dense Mode on this interface	Works	Works	Works
Bidir group	Not supported	Works	Works	Works
Dense group	Works	Suppressed on this interface	Works	Not supported

The ability in IPv4 Cisco IOS Software multicast to manually configure the operational mode of a PIM interface as either **ip pim sparse-mode** or **ip pim dense-mode** was originally provided to allow a router to be configured as a boundary router between a Sparse mode domain and a Dense mode domain. These commands “hardwired” the interface to behave in either a Sparse mode or Dense mode fashion without regard to the actual group mode.

When Auto-RP was introduced, it required that the two Auto-RP multicast groups were always Dense-mode flooded. This resulted in the addition of the **ip pim sparse-dense-mode** command which provided a “dynamic” interface behavior whereby Dense mode groups would be flooded out the interface using Dense mode operation while Sparse mode groups would be forwarded using Sparse mode operation. This allowed the two Auto-RP groups to be Dense mode flooded while groups that had a RP defined for it would be considered as Sparse mode groups and would be forwarded out the interface using Sparse mode forwarding.

The unfortunate side-affect of this approach was that networks that employed Auto-RP required the use of **ip pim sparse-dense-mode** and could therefore “fallback” into Dense mode if the routers lost RP information due to failure of all Candidate RPs for a group. When this occurred, all forwarding state in the routers switched to Dense group mode. This in turn caused the routers to begin Dense mode flooding of all multicast traffic out the **ip pim sparse-dense-mode** interfaces.

Because Cisco IOS Software IPv6 Multicast does not support PIM Dense Mode, there is no need to support multiple PIM interface modes. As a result, Dense mode fallback and Dense mode flooding is not a problem because all non-SSM groups in Cisco IOS Software IPv6 Multicast automatically default to PIM Sparse Mode.

Note: Current industry trends indicate that there is little if any demand for the support of PIM Dense mode in IPv6 Multicast. However, should this trend change Cisco will consider adding support at a later date. If support for PIM Dense mode is added to Cisco IOS Software IPv6 Multicast implementation it is anticipated that the only change to the PIM interface configuration would be to add a special command to identify an interface as being on the Sparse mode – Dense mode domain boundary where Dense mode flooding and proxy Source Registration should take place.

Note: A fix to the problem of IPv4 Dense mode Fallback (and the resultant Dense mode flooding on interfaces configured with **ip pim sparse-dense-mode**) will be available in Cisco IOS Software releases 12.2S (not integrated yet). A new global command, **no ip pim dm-fallback**, will prevent the router from “falling back” into Dense mode should a router experience loss of RP information for a multicast group.

Note: A fix to the problem that prevented Auto-RP from working properly over interfaces configured in **ip pim sparse-mode** was introduced in Cisco IOS Software releases 12.1.13E. Beginning with this release, customers who wish to run Auto-RP and want to avoid the risk of Dense-mode flooding on interfaces configured with **ip pim sparse-dense-mode**, can make use of the new global command, **ip pim autorp listener** to force Auto-RP to work over interfaces configured in **ip pim sparse-mode**.

APPENDIX B—BOOTSTRAP ROUTER EFT IMPLEMENTATION

The BSR [\[17\]](#) EFT Cisco IOS Software has been tested on 6net [\[22\]](#). It allows for most of the functionality of the IETF draft, only full scoping was not fully tested yet.

Configuration Commands

```
[no] ipv6 pim bsr candidate-bsr <address> [<hash-mask-len>] [priority <priority>] [scope <3-15>]
```

Command used to configure a router as a candidate BSR. When configured the router will participate in BSR election. If elected BSR, this router will periodically originate BSR messages advertising the group to rp mappings it has learnt via C-RP Advertisements. If a scope is specified, the BSR will originate BSMs including the group range associated with the scope & accept C-RP announcements only if they are for groups that belong to the given scope.

```
[no] ipv6 pim bsr border
```

Command used to configure a border for all BSM messages of any scope. This command will prevent all BSR messages from being forwarded or accepted on the interface.

```
[no] ip pim bsr candidate-rp <address> [group-list <acl>] [priority <prio>] [bidir] [scope <3-15>]
```

Configures to send pim version 2 candidate RP advertisement to the BSR. The group prefixes defined by named access-list <acl> will also be advertised in association with the RP address. If a group prefix in the access list is denied it will NOT be included in the C-RP advertisement.

If the keyword “bidir” is supplied, the group range will be used for bidirectional shared-tree forwarding otherwise it will be used for PIM-SM forwarding (bidir keyword not yet available in the tested EFT). If the option “priority <prio>” is used, then the router will announce itself to be a candidate RP with priority <prio>. The default for <prio> is 192 and is not NVgened. If scope is specified, then this router will advertise itself as C-RP only to the BSR for the specified scope.

Show Commands

```
show ipv6 pim bsr election
```

Displays currently known BSR state, BSR election & BSM related timers.

```
show ipv6 pim bsr rp-cache
```

Displays the C-RP cache, learnt from unicast C-RP announcements on the elected BSR.

```
show ipv6 pim bsr candidate-rp
```

Displays Candidate RP state on router that are configured as C-RPs. Information on scope, uptime, time until next advertisement is sent etc. is displayed.

```
show ipv6 pim group-mapping [[<group> | < group-name>] | [[group-range/mask] [info-source  
bsr|static|default]]]
```

Enhanced to display group to RP mapping cache learnt from BSR. When a fully qualified group name or group address is supplied it will return the group-mapping that will be used for the particular group. The info-source keyword can be used to limit the ranges displayed to just those enabled by default or via static configuration or learnt via BSR.

Debug Commands

```
debug ipv6 pim bsr
```

Displays debugs specific to BSR protocol operation.

REFERENCES

- [1] IPv6 Multicast Configuration in Cisco IOS Software:
http://www.cisco.com/en/US/tech/tk828/technologies_white_paper09186a0080203f7a.shtml
- [2] IANA Allocated Permanent Multicast addresses: <http://www.iana.org/assignments/ipv6-multicast-addresses>
- [3] Embedding the Rendezvous Point (RP) Address in an IPv6 Multicast Address: <http://www.ietf.org/rfc/rfc3956.txt>
- [4] IPv6 Multicast address allocation study—The 6net WP3 deliverable D3.4.3 available at: <http://www.6net.org>
- [5] Routing IPv6 with IS-IS: <http://www.ietf.org/internet-drafts/draft-ietf-isis-ipv6-05.txt>
- [6] Configuring IPv6 IS-IS:
http://www.cisco.com/en/US/customer/products/sw/iosswrel/ps5187/products_configuration_guide_chapter09186a00801d65f6.html
- [7] M-ISIS: Multi Topology (MT) Routing in IS-IS: <http://www.ietf.org/internet-drafts/draft-ietf-isis-wg-multi-topology-07.txt>
- [8] IS-IS Multi topology Support for IPv6:
http://www.cisco.com/en/US/products/sw/iosswrel/ps5187/products_configuration_guide_chapter09186a00801d65f6.html
http://www.cisco.com/en/US/customer/products/sw/iosswrel/ps5187/products_configuration_guide_chapter09186a00801d65f6.html
- [9] Configuring IPv6 RIP: <http://www.cisco.com/univercd/cc/td/doc/product/software/ios122/122newft/122t/122t2/ipv6/ftip6c.htm - 1018945>
- [10] Implementing OSPF for IPv6:
http://www.cisco.com/univercd/cc/td/doc/product/software/ios122/122newft/122t/122t13/ipv6/ipv6imp/sa_ospf3.htm - wp1076395
- [11] IPv6 Multicast address-family in OSPFv3—deleted from the IETF web site at the moment (not updated <http://www.ietf.org/internet-drafts/draft-mirtorabi-ospfv3-multicast-af-00.txt>)
- [12] Address Family Numbers: <http://www.iana.org/assignments/address-family-numbers>
- [13] Configuring Multiprotocol BGP Extensions for IPv6:
<http://www.cisco.com/univercd/cc/td/doc/product/software/ios122/122newft/122t/122t2/ipv6/ftip6c.htm - 1004556>
- [14] Protocol Independent Multicast—Sparse Mode (PIM-SM):
Protocol Specification (Revised) <http://www.ietf.org/internet-drafts/draft-ietf-pim-sm-v2-new-11.txt>
- [15] PIM upstream detection among multiple addresses—obsoleted and included in the PIM specification: <http://www.ietf.org/internet-drafts/draft-suz-pim-upstream-detection-00.txt>
- [16] IPv6 Multicast Commands—Nidhi Bhaskar—Oct 2002, multicast-v6-commands.txt
- [17] Bootstrap Router (BSR) Mechanism for PIM Sparse Mode: <http://www.ietf.org/internet-drafts/draft-ietf-pim-sm-bsr-04.txt>
- [18] Multicast Listener Discovery Version 2 (MLDv2) for IPv6: <http://www.ietf.org/rfc/rfc3810.txt>
- [19] M6Bone: an experimental IPv6 Multicast network: <http://www.m6bone.net/topology.html>
- [20] Source-Specific Multicast for IP: <http://www.ietf.org/internet-drafts/draft-ietf-ssm-arch-06.txt>
- [21] Source-Specific Protocol Independent Multicast in 232/8: <http://www.ietf.org/internet-drafts/draft-ietf-mboned-ssm232-08.txt>
- [22] 6net—The IPv6 Multicast Testbed—BSR EFT, Gunter Van de Velde

**Corporate Headquarters**

Cisco Systems, Inc.
170 West Tasman Drive
San Jose, CA 95134-1706
USA
www.cisco.com
Tel: 408 526-4000
800 553-NETS (6387)
Fax: 408 526-4100

European Headquarters

Cisco Systems International BV
Haarlerbergpark
Haarlerbergweg 13-19
1101 CH Amsterdam
The Netherlands
www-europe.cisco.com
Tel: 31 0 20 357 1000
Fax: 31 0 20 357 1100

Americas Headquarters

Cisco Systems, Inc.
170 West Tasman Drive
San Jose, CA 95134-1706
USA
www.cisco.com
Tel: 408 526-7660
Fax: 408 527-0883

Asia Pacific Headquarters

Cisco Systems, Inc.
168 Robinson Road
#28-01 Capital Tower
Singapore 068912
www.cisco.com
Tel: +65 6317 7777
Fax: +65 6317 7799

Cisco Systems has more than 200 offices in the following countries and regions. Addresses, phone numbers, and fax numbers are listed on the Cisco Web site at www.cisco.com/go/offices.

Argentina • Australia • Austria • Belgium • Brazil • Bulgaria • Canada • Chile • China PRC • Colombia • Costa Rica • Croatia • Cyprus • Czech Republic • Denmark • Dubai, UAE • Finland • France • Germany • Greece • Hong Kong SAR • Hungary • India • Indonesia • Ireland • Israel • Italy • Japan • Korea • Luxembourg • Malaysia • Mexico • The Netherlands • New Zealand • Norway • Peru • Philippines • Poland • Portugal • Puerto Rico • Romania • Russia • Saudi Arabia • Scotland • Singapore • Slovakia • Slovenia • South Africa • Spain • Sweden • Switzerland • Taiwan • Thailand • Turkey • Ukraine • United Kingdom • United States • Venezuela • Vietnam • Zimbabwe

Copyright © 2004 Cisco Systems, Inc. All rights reserved. CCIP, CCSP, the Cisco *Powered* Network mark, Cisco Unity, Follow Me Browsing, FormShare, and StackWise are trademarks of Cisco Systems, Inc.; Changing the Way We Work, Live, Play, and Learn, and iQuick Study are service marks of Cisco Systems, Inc.; and Aironet, ASIST, BPX, Catalyst, CCDA, CCDP, CCIE, CCNA, CCNP, Cisco, the Cisco Certified Internetwork Expert logo, Cisco IOS, the Cisco IOS logo, Cisco Press, Cisco Systems, Cisco Systems Capital, the Cisco Systems logo, Empowering the Internet Generation, Enterprise/Solver, EtherChannel, EtherSwitch, Fast Step, GigaStack, Internet Quotient, IOS, IP/TV, iQ Expertise, the iQ logo, iQ Net Readiness Scorecard, LightStream, Linksys, MGX, MICA, the Networkers logo, Networking Academy, Network Registrar, *Packet*, PIX, Post-Routing, Pre-Routing, RateMUX, Registrar, ScriptShare, SlideCast, SMARTnet, StrataView Plus, Stratum, SwitchProbe, TeleRouter, The Fastest Way to Increase Your Internet Quotient, TransPath, and VCO are registered trademarks of Cisco Systems, Inc. and/or its affiliates in the United States and certain other countries.

All other trademarks mentioned in this document or Web site are the property of their respective owners. The use of the word partner does not imply a partnership relationship between Cisco and any other company. (0402R)

204032.2_ETMG_AE_12.04