



BGP Next-Hop Processing

This chapter explains how BGP handles next hop addresses for route processing, including notification mechanisms, configuration commands, route policies, and advanced use-cases such as MPLS LSP resolution.

- [BGP next hop notifications, on page 1](#)
- [VRF next hop route policies, on page 7](#)
- [BGP nexthop resolutions over MPLS LSPs with RSVP-TE tunnels, on page 9](#)

BGP next hop notifications

A BGP next hop notification is a dynamic route monitoring mechanism that

- triggers updates to BGP route processing when a change in next hop reachability, connectivity, locality, or IGP metric is detected
- distinguishes between critical and noncritical event types for efficient network response, and
- supports policy-based filtering and batching to minimize route oscillation and ensure stable routing.

Table 1: Feature history table

Feature Name	Release Information	Feature Description
BGP Next Hop Tracking	Release 24.4.1	Introduced in this release on: Fixed Systems (8200 [ASIC: P100], 8700 [ASIC: P100])(select variants only*); Modular Systems (8800 [LC ASIC: P100])(select variants only*) *This feature is supported on: • 88-LC1-36EH+A8:B12 • 88-LC1-12TH24FH-E • 88-LC1-52Y8H-EM • 8212-48FH-M • 8711-32FH-M • 8712-MOD-M

Feature Name	Release Information	Feature Description
BGP Next Hop Tracking	Release 24.3.1	Introduced in this release on: Fixed Systems (8200 [ASIC: Q200, P100], 8700 [ASIC: P100])(select variants only*); Modular Systems (8800 [LC ASIC: Q100, Q200, P100])(select variants only*) The BGP next-hop tracking feature allows refined route resolution, avoiding aggregate routes and oscillation risks by filtering based on prefix length and source protocols, configurable through the nexthop trigger-delay and nexthop route-policy commands. * The BGP next hop tracking functionality is now extended to: • 8212-48FH-M • 8711-32FH-M • 88-LC1-52Y8H-EM • 88-LC1-12TH24FH-E

Event classification and handling

BGP classifies next hop events into two types:

- **Critical events:** These include changes in reachability (reachable or unreachable), connectivity (connected or disconnected), and locality (local or nonlocal). Critical events are processed and sent to BGP immediately to minimize convergence time.
- **Noncritical events:** These are changes in IGP metrics. Noncritical events are collected and sent in batches every three seconds. This batching helps minimize route churn and maintain network stability.

Events that trigger notifications

BGP receives next hop notifications when any of these events occur:

- A next hop becomes reachable or unreachable.
- A next hop becomes connected, unconnected, local, or nonlocal.
- The first hop IP address or interface changes.
- The recursed IGP metric for the next hop changes (e.g., due to network congestion or topology updates).

Configuration options for event handling

Several commands are available to fine-tune how BGP handles next hop notifications:

- **nexthop trigger-delay:** Adjusts the batching interval for noncritical events, per address family, allowing network operators to balance convergence speed with network stability.
- **nexthop route-policy:** Enables policy-based filtering of notifications to control which next hop changes are significant for BGP processing.

Example Scenarios:

- **Critical Event:** If an interface fails and a next hop becomes unreachable, BGP receives an immediate notification and recalculates the best path.
- **Noncritical Event:** When the IGP metric to a next hop increases due to network congestion, BGP receives a noncritical notification and may update its path selection after batching the event.

How BGP next hop notifications work

Summary

The key components involved in the process are:

- **Routing Information Base (RIB):** A database that continuously monitors the status and attributes of all next hops used by BGP.
- **BGP process:** A process that receives notifications from the RIB and updates routing information and neighbor relationships as needed.
- **Network events:** Events that trigger changes, such as interface status or IGP metric adjustments, that require BGP response.

The RIB continuously monitors BGP next hops and notifies BGP of any status changes. BGP then evaluates the event's impact, updates path selection, and recalculates next hop values to ensure correct routing.

Workflow

These stages describe the BGP next hop notifications process:

1. **Monitoring:** The RIB tracks all BGP next hops, continuously monitoring their reachability, connectivity, locality, and associated IGP metrics.
2. **Detection:** The RIB detects a change in next hop status such as an interface going down, a next hop becoming unconnected, or an IGP metric update.
3. **Notification:** Upon detecting a change, the RIB immediately notifies the BGP process about the specific event and affected next hop.
4. **Evaluation:** The BGP process verifies the new state of the next hop and determines whether the event is critical (immediate action needed) or noncritical (batching possible).
5. **Update:** BGP updates its path selection, recalculates outgoing next hop values for all affected routes, and checks or re-establishes neighbor connectivity as required.

Next hop determination (IPv4 and IPv6)

BGP uses a set of rules to determine the next hop address for prefixes exchanged between BGP peers. The determination logic differs depending on whether the prefixes are IPv4 or IPv6 and on the interface and neighbor configurations.

Next hops as IPv6 addresses of peering interfaces

When BGP is used for IPv6 routing, the next hop for IPv6 prefixes can be set to the IPv6 address of the peering interface. This behavior is defined as follows:

- **Assignment:**

The IPv6 address of the peering interface is assigned as the next hop for IPv6 prefixes exchanged over an IPv4 BGP session.

- **Automatic application:**

This assignment is automatically applied when no nexthop policy is configured but either an IPv6 neighbor or an IPv6 update-source interface is present.

- **Default behavior:**

If neither an IPv6 neighbor interface nor an IPv6 update-source interface is configured, the next hop defaults to an IPv4-mapped IPv6 address.

IPv6 prefix transport over IPv4 BGP sessions

BGP supports the transport of IPv6 prefixes over IPv4 sessions. The determination of the next hop in these cases is controlled by the presence or absence of specific interface configurations and policies:

- With IPv6 neighbor or update-source configured:

If an IPv6 neighbor interface or IPv6 update-source interface is configured, BGP uses the IPv6 address of that interface as the next hop.

- Without IPv6 neighbor or update-source:

If neither interface is configured, BGP sets the next hop as an IPv4-mapped IPv6 address.

Table-policy for global IPv6 next hop

By default, BGP switches to a link-local IPv6 address for the next hop when ECMP links change, which may cause transient traffic loss. To maintain consistent load balancing and avoid these issues, configure a BGP table-policy to set a global IPv6 next hop.

Configure a BGP table-policy to set the global IPv6 next hop

Ensure BGP uses a global IPv6 address as the next hop and prevent traffic loss during ECMP path changes.

Before you begin

- Ensure you have access to the BGP route-policy configuration on your device.
- Ensure you have identified the destination prefixes that require consistent load balancing and global IPv6 next hop assignment.

Procedure

Define a new route-policy to set the global IPv6 address as the next hop.

Example:

```
Router(config)# route-policy RESILIENT-HASH-V6
Router(config-rpl)# if destination in (2001:DB8::/32 le 128) then
Router(config-rpl-if)# set load-balance ecmp-consistent
Router(config-rpl-if)# set next-hop ipv6-global
Router(config-rpl-if)# else
Router(config-rpl-else)# pass
Router(config-rpl-else)# endif
Router(config-rpl)# end-policy
```

Route resolution policies

Route resolution policies in BGP provide mechanisms to control how next hops are resolved and which routes are considered valid for advertisement or installation in the routing table. These policies offer fine-grained control over routing decisions and help prevent issues such as route oscillation or unwanted route propagation.

Policy criteria

BGP allows the creation and application of route policies that can specify:

- Prefix length requirements: For example, only allowing next hop routes with a prefix length greater than a configured value.
- Source protocol requirements: Requiring that next hop routes be learned from specific routing protocols (such as Intermediate System to Intermediate System (IS-IS), Open Shortest Path First (OSPF), or connected routes) to ensure routing stability.

Enabling route policy filtering

Enable route policy filtering in BGP using the **nexthop route-policy** command. You can apply this command globally, to specific address families, or at the VRF level, depending on your network design needs.

Scoped IPv4 table walks

A scoped IPv4 table walk is a route lookup mechanism that

- receives next-hop notifications to trigger address family processing
- uses gateway context information to determine which address families share the gateway context, and
- localizes processing to the IPv4 unicast address family table using a next-hop mask.

Identifying the address family from next-hop notifications

When a next-hop notification is received, the system first de-references the gateway context associated with that next hop. It then examines the gateway context to determine which address families are using it.

Gateway context sharing among IPv4 unicast address families

IPv4 unicast address families share the same gateway context, as they are registered with the IPv4 unicast table in the Routing Information Base (RIB). Each next-hop entry includes a mask to show whether it belongs to the IPv4 unicast address family.

Scoped table walk for efficient processing

Whenever a next-hop notification for IPv4 unicast is received from the RIB, the system processes only the global IPv4 unicast table. This scoped table walk ensures that updates or lookups are performed only in the relevant address family table, improving efficiency by avoiding unnecessary processing of unrelated address families.

Address family processing order

When a next-hop notification batch is received, the software reorders the address family processing in this order:

- IPv4 tunnel
- VPNv4 unicast
- IPv4 labeled unicast
- IPv4 unicast
- IPv4 multicast
- IPv6 unicast

This order determines how address family tables are walked based on their numeric value, ensuring efficient routing table updates in response to notifications.

Critical-event thread for BGP next hop processing

The dedicated critical-event thread in the spkr process handles next-hop, BFD, and fast-external-failover (FEF) notifications to help you achieve fast BGP convergence, even during other time-consuming events.

Configuration and commands

BGP provides a variety of configuration and operational commands to help users manage, monitor, and troubleshoot next hop processing. These commands allow you to control next hop handling in BGP updates, gather statistical information, and perform diagnostic analysis of next hop-related events.

Common BGP next hop commands

Show commands:

- **show bgp nexthops:**

Displays statistical information about next hop notifications, processing time, and details about each next hop registered with the RIB. Use this command to monitor the status and performance of next hop processing.

Clear commands:

- **clear bgp nexthop performance-statistics:** Clears the cumulative statistics related to BGP next hop notification processing. Use this command when you want to reset the counters and start a new monitoring interval.

- **clear bgp nexthop registration:** Performs asynchronous registration of the next hop with the RIB. Use this command to reset and re-register next hops in case of inconsistencies or for troubleshooting purposes.

Debug commands:

- **debug bgp nexthop:** Provides diagnostic information on next hop processing.
 - The **out** keyword gives debug information about BGP registration of next hops with the RIB.
 - The **in** keyword displays debug information about next hop notifications received from the RIB.
 - If the **out** keyword is repeated, it displays debug information about next hop notifications sent to the RIB.

Disable next-hop processing on BGP updates

Ensure all BGP updates sent to a neighbor use the local router's address as the next hop, preventing automatic recalculation of next-hop values.

Disabling BGP next-hop processing is required when you want your router to appear as the next hop for all advertised routes to a peer. This is useful in designs where control over routing paths is necessary.

Before you begin

- Obtain your autonomous system (AS) number and the neighbor's IP address.

Procedure

Configure the router to advertise itself as the next hop for all routes sent to the neighbor.

Example:

```
Router(config)# router bgp 120
Router(config-bgp)# neighbor 172.16.40.24
Router(config-bgp-nbr)# remote-as 206
Router(config-bgp-nbr)# address-family ipv4 unicast
Router(config-bgp-nbr-af)# next-hop-self
```

You can also disable next-hop processing for a neighbor group or an address family group, depending on your network design requirements.

VRF next hop route policies

A VRF next hop route policy is a routing control mechanism that

- enables you to configure route policies at the BGP next-hop attach point for individual VRF instances,
- limits notifications delivered to BGP for specific prefixes, and
- provides precise traffic engineering and security compliance for each VRF.

Table 2: Feature History Table

Feature Name	Release Name	Description
Virtual Routing Forwarding Next Hop Routing Policy	Release 7.11.1	You can now enable a route policy at the BGP next-hop attach point to limit notifications delivered to BGP for specific prefixes, which equips you with better control over routing decisions, and allows for precise traffic engineering and security compliance for each VRF instance, and helps establish redundant paths specific to each VRF. The feature introduces these changes: CLI: Modified Command: - The <code>nexthop route-policy</code> command is extended to VRF address-family configuration mode. YANG Data Model - New XPaths for <code>Cisco-IOS-XR-ipv4-bgp-cfg.yang</code> <code>Cisco-IOS-XR-um-router-bgp-cfg</code> (see GitHub, YANG Data Models Navigator)

VRF next hop route policies give network administrators fine-grained control over route advertisement and notification within BGP processes. By assigning route policies to specific VRF address families, you can tailor routing behavior, enhance security, and ensure preferred routing paths for different tenants or services on your network.

Configure a VRF next hop policy

Enable and apply a next hop route policy to a VRF table. This allows you to control which routes are advertised to BGP peers based on prefix and protocol.

Use this task to ensure BGP only learns or advertises specific routes within a VRF.

Before you begin

Decide on the prefixes and protocols you want the route policy to match.

Procedure

Step 1 Define a route policy to match desired prefixes and protocols.

Example:

```
Router(config)# route-policy nh-route-policy
Router(config-rpl)# if destination in (10.1.1.0/24) and protocol in (connected, static) then
Router(config-rpl-if)# drop
Router(config-rpl-if)# endif
Router(config-rpl)# end-policy
Router(config-rpl)# exit
```

Step 2 Enter BGP configuration mode, and apply the next hop route policy to the VRF address family.

Example:

```
Router(config)# router bgp 500
Router(config-bgp)# vrf vrf10
Router(config-bgp-vrf)# address-family ipv4 unicast
Router(config-bgp-vrf-af)# nexthop route-policy nh-route-policy
```

Step 3 Use the **show bgp vrf vrf_name ipv4 unicast** command to verify if the policy is applied.

Example:

```
Router# show bgp vrf vrf1 ipv4 unicast
Fri Jul 7 15:51:16.309 +0530
BGP VRF vrf1, state: Active
BGP Route Distinguisher: 1:1
VRF ID: 0x6000000b
BGP router identifier 10.1.1.1, local AS number 65001
Non-stop routing is enabled
BGP table state: Active
Table ID: 0xe000000b RD version: 1356
BGP table nexthop route policy: nh-route-policy --> This is the same route policy that was configured.
BGP main routing table version 1362
BGP NSR Initial initsync version 1355 (Reached)
BGP NSR/ISSU Sync-Group versions 1362/0

Status codes: s suppressed, d damped, h history, * valid, > best
               i - internal, r RIB-failure, S stale, N Nexthop-discard
Origin codes: i - IGP, e - EGP, ? - incomplete
   Network             Next Hop         Metric  LocPrf  Weight Path
Route Distinguisher: 1:1 (default for vrf vrf1)
Route Distinguisher Version: 1356
*> 10.1.1.0/24         0.0.0.0             0         32768   ?
*> 192.0.2.0/24        10.1.1.1            0         32768   ?
```

BGP nexthop resolutions over MPLS LSPs with RSVP-TE tunnels

A BGP nexthop resolution over MPLS LSPs with RSVP-TE tunnels is a BGP feature that

- resolves BGP nexthops over RSVP-TE tunnels instead of native IP paths

- enforces controlled and predictable traffic steering by forwarding prefixes exclusively through MPLS LSPs, and
- prevents traffic drops caused by the lack of downstream BGP routing information in core networks.

Table 3: Feature History Table

Feature Name	Release Information	Feature Description
BGP nexthop resolution over MPLS LSPs with RSVP-TE tunnels	Release 25.3.1	Introduced in this release on: Centralized Systems (8400 [ASIC: K100])(select variants only*) * This feature is now supported on Cisco 8404-SYS-D routers.

Feature Name	Release Information	Feature Description
BGP nexthop resolution over MPLS LSPs with RSVP-TE tunnels	Release 25.1.1	Introduced in this release on: Fixed Systems (8200 [ASIC: Q200, P100], 8700 [ASIC: P100, K100], 8010 [ASIC: A100]); Centralized Systems (8600 [ASIC: Q200]); Modular Systems (8800 [LC ASIC: Q100, Q200, P100]) You can now prevent BGP prefixes from defaulting to native IP paths, which could lead to traffic drops due to the lack of BGP routing information on downstream core routers, by enforcing BGP nexthop resolution over MPLS LSPs with RSVP-TE tunnels. The feature gives you precise control over traffic steering by defining how BGP resolves nexthops and enabling route policies that consistently forward prefixes over RSVP-TE tunnels. Previously, in core networks, downstream routers without BGP routes caused traffic to default to native IP paths instead of RSVP-TE LSPs, leading to potential drops. The feature introduces these changes: CLI: The next-hop-type is introduced as a filter type in RPL. YANG Data Models: Cisco-IOS-XR-um-route-policy-cfg (see GitHub, YANG Data Models Navigator)

BGP nexthop resolution over MPLS LSPs with RSVP-TE tunnels enables network operators to control BGP nexthop selection using route policies, ensuring prefixes are always forwarded over RSVP-TE engineered paths. This approach avoids traffic drops that can occur when downstream routers lack BGP routing information and traffic otherwise defaults to native IP paths. By integrating filtering based on route-type and next-hop-type, the feature optimizes path selection, prevents congestion, and supports predictable traffic engineering.

Importance in multihomed BGP environments

This feature is critical in multihomed BGP environments, where destination prefixes are reachable via multiple Label Edge Routers (LERs). It ensures that ingress LERs choose nexthops resolving over RSVP-TE LSPs, and not native IP paths, helping keep traffic within engineered tunnels.

Effect on traffic when RSVP-TE tunnels Are unavailable

When no RSVP-TE tunnel is available, BGP marks routes as inaccessible and drops packets, enforcing strict resolution over RSVP-TE tunnels. If a tunnel fails, BGP reroutes traffic through another available RSVP-TE tunnel to maintain continuous traffic flow.

Controlled traffic steering with RSVP-TE tunnels

Suppose a service provider core network has multiple downstream routers, some of which lack complete BGP information. With this feature enabled, BGP will resolve traffic only over RSVP-TE tunnels, which avoids unexpected drops even when native IP paths are present.

Benefits of nexthop resolution with RSVP-TE tunnels

Nexthop resolution over MPLS Label Switched Paths (LSPs) with RSVP-TE tunnels provides several key advantages in scalable network environments:

- Controls traffic steering by directing BGP prefixes over engineered traffic engineering (TE) paths, improving reliability and reducing packet drops.
- Enables flexible integration with multiple transport types and protocols, supporting seamless network expansion and future upgrades.
- Supports thousands of BGP nexthops and prefixes, maintaining high performance in large-scale deployments.

How BGP nexthop resolution over RSVP-TE tunnels works

Summary

The key components involved in the process are:

- BGP process: A process that exchanges routing updates for destination prefixes and determines next-hop addresses.
- Routing Information Base (RIB): A database that monitors next-hop resolution status and notifies BGP of changes, such as next-hop types or availability.
- RSVP-TE tunnels: A tunnel that establishes label-switched paths for traffic forwarding.

BGP uses RSVP-TE tunnels for nexthop resolution by applying route policies and selecting engineered paths. If an RSVP-TE tunnel becomes unavailable, affected routes are marked inaccessible and traffic is dropped until the tunnel is restored.

Workflow

These stages describe the BGP nexthop resolution over RSVP-TE tunnels process.

1. The BGP process receives routing updates for destination prefixes and attempts to resolve each prefix's next-hop address.
2. The RIB monitors the status of next-hop resolution and notifies BGP of any changes, including a change in next-hop type (such as resolution over an RSVP-TE tunnel).

3. BGP evaluates available paths by applying configured route policies and selects those that can be resolved over RSVP-TE tunnels.
4. The router forwards packets based on the best path as determined by BGP, steering traffic over the RSVP-TE tunnels specified.
5. If the RSVP-TE tunnel for a route's next-hop becomes unavailable, BGP marks the route as inaccessible and the router drops traffic for affected prefixes until resolution is restored.

Best practice for configuring BGP nexthop resolution over RSVP-TE tunnels

When configuring BGP nexthop resolution over RSVP-TE tunnels, follow these best practices:

- Ensure that route policies with extended filters are applied to the VPNv4, VPNv6, IPv4, and IPv6 address families.
- Apply the feature to all supported routes within the relevant address families to maintain consistent policy enforcement.

Configure BGP nexthop resolution with RSVP-TE route policy

Limit BGP nexthop resolution to RSVP-TE tunnels by applying a custom route policy.

Before you begin

Ensure RSVP-TE tunnels are established and operational within your network.

Procedure

- Step 1** Define a route policy that permits only nexthops resolved using MPLS-TE.

Example:

```
Router(config)#route-policy ROUTE-RESOLUTION
Router(config-rpl)#if protocol is isis 100 and route-type is level-1 then
Router(config-rpl-if)#pass
Router(config-rpl-if)#elseif protocol is isis 100 and route-type is level-2 and next-hop-type is
mpls-te then
Router(config-rpl-elseif)#pass
Router(config-rpl-elseif)#else
Router(config-rpl-else)#drop
Router(config-rpl-else)#endif
Router(config-rpl)#end-policy
Router(config)#commit
```

- Step 2** Apply the route policy to the BGP address-family configuration.

Example:

```
Router(config)# router bgp 100
Router(config-bgp)# address-family ipv4 unicast
Router(config-bgp-af)# nexthop route-policy ROUTE-RESOLUTION
Router(config-bgp-af)# commit
Router(config-bgp)# address-family ipv6 unicast
Router(config-bgp-af)# nexthop route-policy ROUTE-RESOLUTION
Router(config-bgp-af)# commit
```

Step 3 Use the **show bgp ipv4 unicast nexthops** command to verify if BGP resolves the nexthop over RSVP-TE tunnels.

Example:

```
Router# show bgp ipv4 unicast nexthops 209.165.200.225
Tue Mar 18 04:47:32.759 UTC
Nexthop: 129.134.99.7
  VRF: default
  Nexthop ID: 0x60000004, Version: 0
  Nexthop Flags: 0x00000000
  Nexthop Handle: 0x7f8d5fab66a8
  Tree Nexthop Handle: 0x7f8d5fab66a8

RIB Related Information:
Firsthop interface handle 0x7800004c
  Gateway TBL Id: 0xe0000000 Gateway Flags: 0x00000080
  Gateway Handle: 0x7f8d5d8b9d10
  Gateway: reachable, non-Connected route, prefix length 32
  Resolving Route: 129.134.99.7/32 (isis 1)
  Paths: 0
  RIB Nexthop Route Type: ISIS level-2
  RIB Nexthop Path Type: MPLS-TE
  RIB Nexthop ID: 0x0
  Nexthop sync slot: 15
  Status: [Reachable][Not Connected][Not Local]
  Metric: 30
  ORR afi bits: 0x0
  Inactive Tables: [IPv6 Unicast]
  Registration: Asynchronous, Completed: 00:27:39
  Events: Critical (3)/Non-critical (0)
  Last Received: 00:01:21 (Critical)
  Last gw update: (Crit-notif) 00:01:21(rib)
  Reference Count: 10

  Reachable Notifications: 1 (last at Mar 18 04:19:54.340)
  Unreachable Notifications: 0
  Metric Increase Notifications: 0
  Metric Decrease Notifications: 2
  Nexthop find: 8
  Most Recent Events:
    Time           EventType    Metric  Ifhandle    RouteType    PathType
    Mar 18 04:19:54.340 Reachable    30      0x7800004c  ISIS level-2  MPLS-TE
    Mar 18 04:44:47.974 Reachable    30      0x78000014  ISIS level-2  any
    Mar 18 04:46:12.411 Reachable    30      0x7800004c  ISIS level-2  MPLS-TE

Prefix Related Information
Active Tables: [IPv4 Unicast]
Metrics: [0x1e]
Reference Counts: [10]
Encapsulations: []
Interface Handle: 0x0
Attr ref-count: 13
```

Review the output for PathType MPLS-TE, indicating successful resolution over RSVP-TE tunnels.

Step 4 Use the **show bgp ipv4 unicast** command to verify BGP routing information and nexthop reachability.

Example:

```
Router# show bgp ipv4 unicast 209.165.201.1/27
Thu Feb 13 12:14:35.626 UTC
BGP routing table entry for 209.165.201.1/27
Versions:
  Process          bRIB/RIB    SendTblVer
  Speaker          38992497    38992497
```

```

Last Modified: Feb 13 12:14:25.298 for 00:00:10
Paths: (1 available, no best path)
  Not advertised to any peer
  Path #1: Received by speaker 0
  Not advertised to any peer
  64800, (received & used)
    10.1.9.7 (inaccessible) from 10.3.9.7 (10.4.9.7)
      Origin IGP, localpref 100, valid, internal
      Received Path ID 0, Local Path ID 0, version 0
      Community: 65530:50700

```

If no RSVP-TE tunnel exists, the nexthop appears as `inaccessible`.

Step 5 Use the `show rib ipv4 unicast next-hop` command to confirm if the RIB is forwarding the route over MPLS-TE tunnels.

Example:

```

Router#show rib ipv4 unicast next-hop 209.165.201.1/27

Tue Mar 18 04:47:33.806 UTC

Firsthop prefix: 209.165.201.1/27
Flags: allow default, recurse
Ext flags: 0x1 (all_path_mpls_te)
Damped counter: 0
Damp algo hits: 3
Last event occurred Mar 18 04:46:12.409, 00:01:21 ago; version 4

Registered clients:
  te_control/node0_RP0_CPU0 created Mar 18 04:19:23.479, 00:28:10 ago
    read last notification at Mar 18 04:46:12.411, 00:01:21 ago
    reference count is 1

Destination paths:
  209.165.201.1 - S-AR1-DR1-1-1
  209.165.201.1 - S-AR1-DR1-2-1
  209.165.201.1 - S-AR1-DR1-3-1
  209.165.201.1 - S-AR1-DR1-4-1
  209.165.201.1 - S-AR1-DR1-5-1
  209.165.201.1 - S-AR1-DR1-6-1
  209.165.201.1 - S-AR1-DR1-7-1
  209.165.201.1 - SF-AR1-DR1-1-1

Resolving route: 209.165.201.1/27 known via "isis 1"
Metric computed: 30

```

If `Ext flags: 0x1 (all_path_mpls_te)` is present, the route is using only MPLS-TE tunnels.

