



7705 Service Aggregation Router Gen 2

Release 26.7.R1

Segment Routing and PCE User Guide

3HE 29569 AAAB TQZZA 01
Edition: 01
July 2026

Nokia is committed to diversity and inclusion. We are continuously reviewing our customer documentation and consulting with standards bodies to ensure that terminology is inclusive and aligned with the industry. Our future customer documentation will be updated accordingly.

This document includes Nokia proprietary and confidential information, which may not be distributed or disclosed to any third parties without the prior written consent of Nokia.

This document is intended for use by Nokia's customers ("You"/"Your") in connection with a product purchased or licensed from any company within Nokia Group of Companies. Use this document as agreed. You agree to notify Nokia of any errors you may find in this document; however, should you elect to use this document for any purpose(s) for which it is not intended, You understand and warrant that any determinations You may make or actions You may take will be based upon Your independent judgment and analysis of the content of this document.

Nokia reserves the right to make changes to this document without notice. At all times, the controlling version is the one available on Nokia's site.

No part of this document may be modified.

NO WARRANTY OF ANY KIND, EITHER EXPRESS OR IMPLIED, INCLUDING BUT NOT LIMITED TO ANY WARRANTY OF AVAILABILITY, ACCURACY, RELIABILITY, TITLE, NON-INFRINGEMENT, MERCHANTABILITY OR FITNESS FOR A PARTICULAR PURPOSE, IS MADE IN RELATION TO THE CONTENT OF THIS DOCUMENT. IN NO EVENT WILL NOKIA BE LIABLE FOR ANY DAMAGES, INCLUDING BUT NOT LIMITED TO SPECIAL, DIRECT, INDIRECT, INCIDENTAL OR CONSEQUENTIAL OR ANY LOSSES, SUCH AS BUT NOT LIMITED TO LOSS OF PROFIT, REVENUE, BUSINESS INTERRUPTION, BUSINESS OPPORTUNITY OR DATA THAT MAY ARISE FROM THE USE OF THIS DOCUMENT OR THE INFORMATION IN IT, EVEN IN THE CASE OF ERRORS IN OR OMISSIONS FROM THIS DOCUMENT OR ITS CONTENT.

Copyright and trademark: Nokia is a registered trademark of Nokia Corporation. Other product names mentioned in this document may be trademarks of their respective owners.

© 2026 Nokia.

Table of contents

List of tables.....	8
List of figures.....	9
1 Getting started.....	11
1.1 About this guide.....	11
1.2 Platforms and terminology.....	11
1.3 Conventions.....	12
1.3.1 Precautionary and information messages.....	12
1.3.2 Options or substeps in procedures and sequential workflows.....	12
2 Segment routing with MPLS data plane (SR-MPLS).....	14
2.1 Segment routing in shortest path forwarding.....	14
2.1.1 Configuring segment routing in shortest path.....	14
2.1.2 Configuring single shared loopback SR SID.....	18
2.1.3 Segment routing shortest path forwarding with IS-IS.....	20
2.1.3.1 IS-IS control protocol changes.....	20
2.1.3.2 Segment routing multitopology considerations.....	23
2.1.3.3 Announcing ELC, MSD-ERLD, and MSD-BMI with IS-IS.....	27
2.1.3.4 EL for IS-IS segment routing.....	28
2.1.3.5 IPv6 segment routing using MPLS encapsulation.....	28
2.1.3.6 Segment routing mapping server function for IPv4 prefixes.....	29
2.1.4 Segment routing shortest path forwarding with OSPF.....	32
2.1.4.1 OSPFv2 control protocol changes.....	32
2.1.4.2 OSPFv3 control protocol changes.....	33
2.1.4.3 Announcing ELC, MSD-ERLD, and MSD-BMI with OSPF.....	34
2.1.4.4 EL for OSPF segment routing.....	34
2.1.4.5 Segment routing mapping server for IPv4 prefixes.....	35
2.1.5 Segment routing with BGP.....	37
2.1.6 Segment routing operational procedures.....	39
2.1.6.1 Prefix advertisement and resolution.....	39
2.1.6.2 Error and resource exhaustion handling.....	40
2.1.7 Segment routing tunnel management.....	45
2.1.7.1 Tunnel MTU determination.....	46

2.1.8	Segment routing local block.....	47
2.1.8.1	Bundling adjacencies in adjacency sets.....	48
2.1.9	Loop-free alternates.....	50
2.1.9.1	Remote LFA with segment routing.....	51
2.1.9.2	Topology-independent LFA.....	54
2.1.9.3	Node protection support in TI-LFA and remote LFA.....	60
2.1.9.4	LFA policies.....	65
2.1.9.5	LFA protection using a segment routing backup node SID.....	80
2.1.9.6	Multihomed prefix LFA extensions in SR-OSPF.....	86
2.1.9.7	Multihomed prefix LFA extensions in SR IS-IS.....	86
2.1.9.8	LFA solution across IGP area or instance boundary using SR repair tunnel in SR-OSPF.....	87
2.1.9.9	LFA solution across IGP area or instance boundary using SR repair tunnel in SR IS-IS.....	91
2.1.10	Segment routing datapath support.....	92
2.1.10.1	Hash label and EL support.....	94
2.1.10.2	TTL or hop-limit field handling.....	94
2.1.11	BGP shortcuts using segment routing tunnels.....	95
2.1.12	BGP labeled route resolution using segment routing tunnels.....	95
2.1.13	Service packet forwarding with segment routing.....	96
2.1.14	Mirror services.....	96
2.1.15	Class-based forwarding for SR-ISIS over RSVP-TE LSPs.....	97
2.1.16	Segment routing traffic statistics.....	98
2.1.17	Configuring BGP-based services for flexible algorithms.....	99
2.2	Establishing segment routing TE LSPs.....	102
2.2.1	SR-TE MPLS support.....	103
2.2.2	SR-TE LSP path computation.....	104
2.2.3	SR-TE LSP path computation using hop-to-label translation.....	105
2.2.4	SR-TE LSP path computation using local CSPF.....	105
2.2.4.1	Extending MPLS and TE database CSPF support to SR-TE LSP.....	106
2.2.4.2	SR-TE specific TE-DB changes.....	107
2.2.4.3	SR-TE LSP and auto-LSP-specific CSPF changes.....	107
2.2.4.4	Delay metric.....	113
2.2.4.5	Ad hoc SR-TE LSP reoptimization on receipt of IGP link events.....	115
2.2.5	SR-TE LSP paths using explicit SIDs.....	115
2.2.6	SR-MPLS IGP shortcuts over SR-TE LSP.....	116
2.2.7	SR-TE LSP protection.....	118

2.2.7.1	Local protection.....	120
2.2.7.2	End-to-end protection.....	121
2.2.8	Static route resolution using SR-TE LSP.....	121
2.2.9	BGP shortcuts using SR-TE LSP.....	122
2.2.10	BGP labeled route resolution using SR-TE LSP.....	122
2.2.11	Service packet forwarding using SR-TE LSP.....	122
2.2.12	Datapath support.....	123
2.2.12.1	SR-TE LSP metric and MTU settings.....	125
2.2.12.2	LSR hashing on SR-TE LSPs.....	126
2.2.13	SR-TE Auto-LSP.....	127
2.2.13.1	Feature configuration.....	128
2.2.13.2	Automatic creation of an SR-TE mesh LSP.....	128
2.2.13.3	Automatic creation of an SR-TE one-hop LSP.....	129
2.2.13.4	Automatic creation of an on-demand SR-TE LSP.....	130
2.2.13.5	Forwarding contexts supported with SR-TE auto-LSP.....	133
2.2.14	Allocation and binding of labels to SR-TE LSPs.....	133
2.2.15	SR-TE LSP traffic statistics.....	134
2.2.15.1	Rate statistics.....	134
2.2.16	SR-TE label stack checks.....	134
2.2.16.1	SR-TE label stack check for services and shortcuts.....	134
2.2.16.2	Control plane handling of egress label stack limitations.....	136
2.2.16.3	Flexible SR-TE label stack allocation for services.....	139
2.2.17	IPv6 traffic engineering.....	140
2.2.17.1	Global configuration.....	141
2.2.17.2	IS-IS configuration.....	142
2.2.17.3	MPLS configuration.....	142
2.2.17.4	IS-IS, BGP-LS, and TE database extensions.....	143
2.2.17.5	IS-IS IPv4/IPv6 SR-TE and IPv4 RSVP-TE feature behavior.....	149
2.2.17.6	IPv6 SR-TE LSP support in MPLS.....	153
2.2.18	OSPF link TE attribute reuse.....	156
2.2.18.1	OSPF application-specific TE link attributes.....	156
2.2.19	Configuring and operating SR-TE.....	158
2.2.19.1	SR-TE configuration prerequisites.....	158
2.2.19.2	SR-TE LSP configuration overview.....	160
2.2.19.3	Configuring SR-TE LSP label stack size.....	160
2.2.19.4	Configuring adjacency SID parameters.....	160

2.2.19.5	Configuring a mesh of SR-TE auto-LSPs.....	161
2.2.20	EL on SR-TE LSPs.....	170
2.3	Segment routing policies.....	171
2.3.1	Statically-configured segment routing policies.....	172
2.3.2	BGP-signaled SR policies.....	174
2.3.3	Segment routing policy path selection and tie-breaking.....	174
2.3.4	Resolving BGP routes to segment routing policy tunnels.....	176
2.3.4.1	Resolving unlabeled IPv4 BGP routes to segment routing policy tunnels.....	176
2.3.4.2	Resolving unlabeled IPv6 BGP routes to segment routing policy tunnels.....	177
2.3.4.3	Resolving label-IPv4 BGP routes to segment routing policy tunnels.....	178
2.3.4.4	Resolving label-IPv6 BGP routes to segment routing policy tunnels.....	179
2.3.4.5	Resolving EVPN-MPLS routes to segment routing policy tunnels.....	181
2.3.4.6	VPRN auto-bind-tunnel using segment routing policy tunnels.....	181
2.3.5	Traffic statistics for segment routing policies.....	182
3	Standards and protocol support.....	184
3.1	Bidirectional Forwarding Detection (BFD).....	184
3.2	Border Gateway Protocol (BGP).....	184
3.3	Bridging and management.....	185
3.4	Certificate management.....	186
3.5	Ethernet.....	186
3.6	Ethernet VPN (EVPN).....	186
3.7	gRPC Remote Procedure Calls (gRPC).....	187
3.8	Intermediate System to Intermediate System (IS-IS).....	187
3.9	Internet Protocol (IP) general.....	188
3.10	Internet Protocol (IP) multicast.....	189
3.11	Internet Protocol (IP) version 4.....	190
3.12	Internet Protocol (IP) version 6.....	190
3.13	Internet Protocol Security (IPsec).....	191
3.14	Label Distribution Protocol (LDP).....	192
3.15	Multiprotocol Label Switching (MPLS).....	193
3.16	Network Address Translation (NAT).....	193
3.17	Network Configuration Protocol (NETCONF).....	193
3.18	Open Shortest Path First (OSPF).....	193
3.19	Path Computation Element Protocol (PCEP).....	194
3.20	Pseudowire (PW).....	195

3.21	Quality of Service (QoS).....	196
3.22	Remote Authentication Dial In User Service (RADIUS).....	196
3.23	Resource Reservation Protocol - Traffic Engineering (RSVP-TE).....	196
3.24	Routing Information Protocol (RIP).....	196
3.25	Segment Routing (SR).....	197
3.26	Simple Network Management Protocol (SNMP).....	197
3.27	Storage sanitization.....	199
3.28	Timing.....	199
3.29	Two-Way Active Measurement Protocol (TWAMP).....	199
3.30	Virtual Private LAN Service (VPLS).....	200
3.31	Yet Another Next Generation (YANG).....	200

List of tables

Table 1: Platforms and terminology.....	11
Table 2: Outcome of LFA policy with protection-type=node.....	66
Table 3: Outcome of LFA policy with protection-type=link.....	66
Table 4: Handling of duplicate SIDs.....	83
Table 5: OSPF Prefix SID sub-TLV main fields.....	84
Table 6: OSPF Prefix SID sub-TLV flags.....	85
Table 7: Datapath support.....	92
Table 8: Values of the frr-overhead parameter.....	135
Table 9: Label stack egress IOM restrictions on FP-based hardware for IPVPN and EVPN services.....	136
Table 10: Maximum available transport labels for IP shortcuts and spoke SDP services.....	138
Table 11: Egress label stack limits for BGP services based on dynamic-egress-label-limit.....	139
Table 12: Legacy link TE TLV support in TE-DB and BGP-LS.....	146
Table 13: Additional link TE TLV support in TE-DB and BGP-LS.....	146
Table 14: Legacy Link TE TLV support in TE-DB and BGP-LS.....	148
Table 15: Additional Link TE TLV support in TE-DB and BGP-LS.....	148
Table 16: Details of link TE advertisement methods.....	150
Table 17: Nokia support for ASLA extended link TLV encoding.....	157
Table 18: Configuration considerations for TE Opaque LSAs.....	158

List of figures

Figure 1: Packet label encapsulation using segment routing tunnel.....	15
Figure 2: Programming multiple tunnels to the same destination.....	41
Figure 3: Handling of the same prefix and SID in different IS-IS instances.....	44
Figure 4: Example of remote LFA topology.....	51
Figure 5: Remote LFA next hop in segment routing.....	53
Figure 6: Selecting link-protect TI-LFA backup path.....	56
Figure 7: TI-LFA backup path via a pseudo-node.....	58
Figure 8: Parallel adjacencies between P and Q nodes.....	59
Figure 9: Application of the TI-LFA algorithm for node protection.....	61
Figure 10: Application of the remote LFA algorithm for node protection.....	63
Figure 11: Application of LFA policy to RLFA and TI-LFA.....	71
Figure 12: Label stack for remote LFA in ring topology.....	80
Figure 13: Backup ABR node SID.....	82
Figure 14: OSPF Prefix SID sub-TLV.....	84
Figure 15: OSPF Prefix SID sub-TLV flags.....	85
Figure 16: Application of MHP LFA to SR-OSPF tunnel of external prefix.....	87
Figure 17: Application of MHP LFA with repair tunnel to SR-OSPF tunnel of external or anycast prefix.....	90
Figure 18: Transport label stack in shortest path forwarding with segment routing.....	93
Figure 19: Label stack reduction in a 3-tier ring topology.....	112
Figure 20: Label stack reduction in the presence of ECMP paths.....	113
Figure 21: SR-TE LSP label stack programming.....	125

Figure 22: VPRN example of an on-demand SR-TE LSP.....	130
Figure 23: Attribute mapping per application.....	153
Figure 24: Multilevel IS-IS topology in the NSP GUI.....	161
Figure 25: Network example with two SR policies.....	172
Figure 26: Relationship between SR policies and paths.....	172

1 Getting started

1.1 About this guide

This guide describes the Nokia SR OS segment routing and PCE functionality.

This guide is organized into functional chapters and provides concepts and descriptions of the implementation flow, as well as Command Line Interface (CLI) syntax and command usage.

Unless otherwise indicated, the topics and commands described in this guide apply only to the 7705 SAR Gen 2 platforms listed in [Platforms and terminology](#).

Command outputs shown in this guide are examples only; actual displays may differ depending on supported functionality and user configuration.



Note: Unless otherwise indicated, CLI commands, contexts, and configuration examples in this guide apply for both the classic CLI and the MD-CLI.

The SR OS CLI trees and command descriptions can be found in the following guides:

- *7705 SAR Gen 2 Classic CLI Command Reference Guide*
- *7705 SAR Gen 2 Clear, Monitor, Show, Tools CLI Command Reference Guide* (for both the MD-CLI and classic CLI)
- *7705 SAR Gen 2 MD-CLI Command Reference Guide*



Note: This guide generically covers Release 26.x.Rx content and may contain some content that will be released in later maintenance loads. See the *SR OS R26.x.Rx Software Release Notes*, part number 3HE 29176 000x TQZZA, for information about features supported in each load of the Release 26.x.Rx software. For a list of features and CLI commands that are present in SR OS but not supported on the 7705 SAR Gen 2 platforms, see "SR OS Features not Supported on SAR Gen 2" in the *SR OS R26.x.Rx Software Release Notes*.

1.2 Platforms and terminology



Note:

Unless explicitly noted otherwise, this guide uses the terminology defined in the following table to collectively designate the specified platforms.

Table 1: Platforms and terminology

Platform	Collective platform designation
7705 SAR-Hx	7705 SAR Gen 2
7705 SAR-Mx	
7705 SAR-1	

1.3 Conventions

This section describes the general conventions used in this guide.

1.3.1 Precautionary and information messages

The following information symbols are used in the documentation.



DANGER: Danger warns that the described activity or situation may result in serious personal injury or death. An electric shock hazard could exist. Before you begin work on this equipment, be aware of hazards involving electrical circuitry, be familiar with networking environments, and implement accident prevention procedures.



WARNING: Warning indicates that the described activity or situation may, or will, cause equipment damage, serious performance problems, or loss of data.



Caution: Caution indicates that the described activity or situation may reduce your component or system performance.



Note: Note provides additional operational information.



Tip: Tip provides suggestions for use or best practices.

1.3.2 Options or substeps in procedures and sequential workflows

Options in a procedure or a sequential workflow are indicated by a bulleted list. In the following example, at step 1, the user must perform the described action. At step 2, the user must perform one of the listed options to complete the step.

Example: Options in a procedure

1. User must perform this step.
2. This step offers three options. User must perform one option to complete this step.
 - This is one option.
 - This is another option.
 - This is yet another option.

Substeps in a procedure or a sequential workflow are indicated by letters. In the following example, at step 1, the user must perform the described action. At step 2, the user must perform two substeps (a. and b.) to complete the step.

Example: Substeps in a procedure

1. User must perform this step.
2. User must perform all substeps to complete this action.
 - a. This is one substep.

- b.** This is another substep.

Nested substeps within a procedure or a sequential workflow are indicated by roman numerals. In the following example, at step 1, the user must perform the described action. At step 2, the user must perform two substeps (a. and b.) to complete the step. At substep b, the user must perform two additional substeps (i. and ii.) to complete the step.

Example: Nested substeps in a procedure

- 1.** User must perform this step.
- 2.** User must perform all substeps to complete this action.
 - a.** This is one substep.
 - b.** User must perform all nested substeps to complete this action.
 - i.** This is a nested substep.
 - ii.** This is another nested substep.

2 Segment routing with MPLS data plane (SR-MPLS)

This section describes:

- Segment Routing (SR) in shortest path forwarding
- SR with Traffic Engineering (SR-TE)
- SR policies

2.1 Segment routing in shortest path forwarding

Segment routing provides support for shortest path routing and source routing using abstract segments for IS-IS and OSPF protocols. A segment can represent a local prefix of a node, a specific adjacency of the node (interface or next hop), a service context, or a specific explicit path over the network. For each segment, the IGP advertises a Segment ID (SID).

When segment routing is used together with the MPLS data plane, the SID is a standard MPLS label. A router forwarding a packet using segment routing pushes one or more MPLS labels.

Segment routing using MPLS labels can be used in both shortest path routing applications and in traffic engineering (TE) applications.

When a received IPv4 or IPv6 prefix SID is resolved, the Segment Routing module programs the Incoming Label Map (ILM) with a swap operation and programs the LTN with a push operation, both of which point to the primary or Loop-Free Alternate (LFA) Next-Hop Label to Forwarding Entry (NHLFE). An IPv4 or IPv6 tunnel to the prefix destination is also added to the TTM and can be used by shortcut applications and Layer 2 and Layer 3 services.

Segment routing provides the remote LFA feature, which expands the coverage of LFA by computing and automatically programming SR tunnels that are used as backup next hops. The SR shortcut tunnels terminate on a remote alternate node that provides loop-free forwarding for packets with resolved prefixes. When the **loopfree-alternates** option is enabled in an IS-IS or OSPF instance, SR tunnels are protected with an LFA backup next hop. If the prefix of a specific SR tunnel is not protected by the base LFA, the remote LFA automatically computes a backup next hop using an SR tunnel if the **remote-lfa** option is also enabled in the IGP instance.

2.1.1 Configuring segment routing in shortest path

Segment routing in an IGP routing instance is enabled using the sequence of commands described in this section.

First, the user configures the global label block, known as the Segment Routing Global Block (SRGB), which is reserved for assigning labels to segment routing prefix SIDs originated by this router. The label range is derived from the system dynamic label range and is not instantiated by default. The range is configured as follows.

```
config>router>mpls-labels>sr-labels start start-value end end-value
```

Next, the user enables the context to configure segment routing parameters within an IGP instance.

```
config>router>isis>segment-routing
config>router>ospf>segment-routing
```

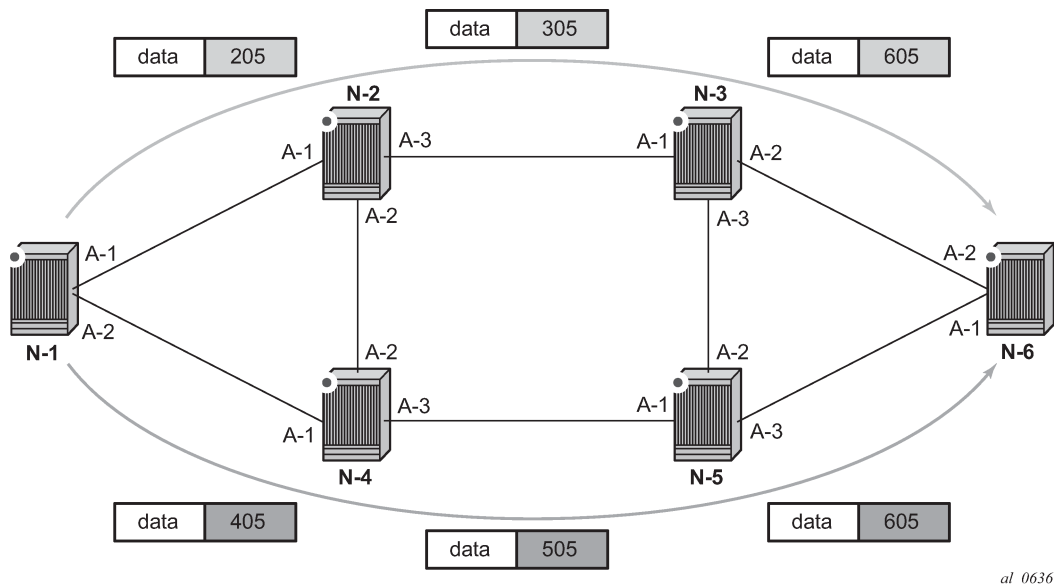
The key parameter is the configuration of the prefix SID index range and the offset label value that this IGP instance uses. Because each prefix SID represents a network global IP address, the SID index for a prefix must be unique network-wide. Thus, all routers in the network are expected to configure and advertise the same prefix SID index range for an IGP instance. However, the label value used by each router to represent this prefix, which is the label programmed in the ILM, can be local to that router by the use of an offset label, referred to as a start label. The relationship between the labels and SIDs is as follows:

Local label (prefix SID) = start label + {SID index}

The label operation in the network is similar to LDP when operating in the independent label distribution mode (RFC 5036), with the difference that the label value used to forward a packet to each downstream router is computed by the upstream router based on the advertised prefix SID index using the above formula.

The following figure shows an example of a router advertising its loopback address and the resulting packet label encapsulation throughout the network.

Figure 1: Packet label encapsulation using segment routing tunnel



Router N-6 advertises loopback 10.10.10.1/32 with a prefix index of 5. Routers N-1 to N-6 are configured with the same SID index range of [1,100] and an offset label of 100 to 600 respectively. The following are the actual label values programmed by each router for the prefix of PE2.

- N-6 has a start label value of 600 and programs an ILM with label 605.
- N-3 has a start label of 300 and swaps incoming label 305 to label 605.
- N-2 has a start label of 200 and swaps incoming label 205 to label 305.

Similar operations are performed by N-4 and N-5 for the bottom path.

N-1 has an SR tunnel to N-6 with two ECMP paths. It pushes label 205 when forwarding an IP or service packet to N-6 via downstream next-hop N-2 and pushes label 405 when forwarding via downstream next-hop N-4.

The CLI syntax for configuring the prefix SID index range and offset label value for an IGP instance are as follows:

```
config>router>isis>segment-routing>prefix-sid-range {global | start-label label-value max-index index-value}
config>router>ospf>segment-routing>prefix-sid-range {global | start-label label-value max-index index-value}
```

There are two mutually-exclusive modes of operation for the prefix SID range on the router. In the global mode of operation, the user configures the global value and this IGP instance takes the start label value as the lowest label value in the SRGB and the prefix SID index range size equal to the range size of the SRGB. After one IGP instance selects the **global** option for the prefix SID range, all IGP instances on the system are restricted to the same **global**.

The user must shut down segment routing context and delete the **prefix-sid-range** command in all IGP instances to change the SRGB. After the SRGB is changed, the user must re-enter the **prefix-sid-range** command. The SRGB range change fails if an already allocated SID index or label goes out of range.

In the per-instance mode of operation, the user partitions the SRGB into non-overlapping subranges among the IGP instances. The user configures a subset of the SRGB by specifying the start label value and the prefix SID index range size. All resulting net label values (start label + index) must be within the SRGB or the configuration fails. Furthermore, the code checks for overlaps of the resulting net label value range across IGP instances and strictly enforces values that do not overlap between ranges.

The user must shut down the segment routing context of an IGP instance to change the SID index or label range of that IGP instance using the **prefix-sid-range** command. Any range change fails if an already allocated SID index or label goes out of range.

The user can change the SRGB at any time as long as it does not reduce the current per-IGP instance SID index or label range defined with the **prefix-sid-range**. Otherwise, the user must shut down the segment routing context of the IGP instance, then delete and reconfigure the **prefix-sid-range** command.

Finally, the user brings up segment routing on that IGP instance by shutting down the context:

```
config>router>isis>segment-routing>no shutdown
config>router>ospf>segment-routing>no shutdown
```

This command fails if the user has not previously enabled the **router-capability** option in the IGP instance. Segment routing must be advertised to all routers in a domain so that routers that support the capability only program the node SID in the datapath toward neighbors that also support it.

```
config>router>isis>advertise-router-capability {area | as}
config>router>ospf>advertise-router-capability {link | area | as}
```

The IGP segment routing extensions are area-scoped. The user must configure the flooding scope as **area** in OSPF and as **area** or **as** in IS-IS.

Next, the user uses one of the following commands to assign a node SID index or label to the prefix, representing the primary address of a network interface of type **system** or **loopback**. A separate SID value can be configured for each IPv4 and IPv6 primary address of the interface.

```
config>router>isis>interface>ipv4-node-sid index value
config>router>ospf>area>interface>node-sid index value
```

```

config>router>ospf3>area>interface>node-sid index value
config>router>isis>interface>ipv4-node-sid label value
config>router>ospf>area>interface>node-sid label value
config>router>ospf3>area>interface>node-sid label value
config>router>isis>interface>ipv6-node-sid index value
config>router>isis>interface>ipv6-node-sid label value

```

The secondary address of an IPv4 interface cannot be assigned a node SID index and does not inherit the SID of the primary IPv4 address. The same applies to the non-primary IPv6 addresses of an interface.

In IS-IS, an interface inherits the configured IPv4 or IPv6 node SID value in any level the interface participates in (Level 1, Level 2, or both).

In OSPFv2 and OSPFv3, the node SID is configured in the primary area but is inherited in any other area in which the interface is added as secondary.

The preceding commands fail if the network interface is not of type **system** or **loopback**, or if the interface is defined in an IES or a VPRN context. Assigning the same SID index or label value to the same interface in two different IGP instances is not allowed within the same node.

For OSPF, the protocol version number and the instance number dictate if the node-SID index or label is for an IPv4 or IPv6 address of the interface. Specifically, the support of address families in OSPF is as follows:

- for ospfv2, always IPv4 only
- for ospfv3, instance 0..31, ipv6 only
- for ospfv3, instance 64..95, ipv4 only

The value of the label or index SID is taken from the range configured for this IGP instance. When using the global mode of operation, a new segment routing module checks that the same index or label value is not assigned to more than one loopback interface address. When using the per-instance mode of operation, this check is not required because the index and the label ranges of the various IGP instances are not allowed to overlap.

For an individual adjacency, values for the label may be provisioned for an IS-IS or OSPF interface. If they are not provisioned, they are dynamically allocated by the system from the dynamic label range. The following CLI commands are used:

```

config>router>isis>interface
[no] ipv4-adjacency-sid label value
[no] ipv6-adjacency-sid label value

config>router>ospf>area>interface
[no] adjacency-sid label value

```

The *value* must correspond to a label in a reserved label block in provisioned mode referred to by the **srlb** command (see [Segment routing local block](#) for more details of SRLBs).

A static label *value* for an adjacency SID is persistent. Therefore, the P-bit of the Flags field in the Adjacency-SID TLV advertised in the IGP is set to 1.

By default, a dynamic adjacency SID is advertised for an interface. However, if a static adjacency SID value is configured, then the dynamic adjacency SID is deleted and only the static adjacency SID used. Changing an adjacency SID from dynamic (for example, **no adjacency-sid**) to static, or the other way around, may result in traffic being dropped as the ILM is reprogrammed.

For a provisioned adjacency SID of an interface, a backup is calculated similar to a regular adjacency SID when **sid-protection** is enabled for that interface.

Provisioned adjacency SIDs are only supported on point-to-point interfaces.

2.1.2 Configuring single shared loopback SR SID

When configuring an IPv4 or IPv6 SR SID for OSPF or IS-IS instances, the single shared SID for loopback or system interfaces can be enabled by the routing protocol independent **sr-mpls>prefix-sids** command. One or more IGP protocol instances can have a unique **sr-mpls>prefix-sids** configured and share this interface SID for an interface. This enhancement relaxes the otherwise imposed SID uniqueness for a loopback or system interface across all configured routing instances on a device.

It is possible to configure the **sr-mpls>prefix-sids** by label or index. The global **prefix-sid-range** must be configured in the routing instance when the **sr-mpls>prefix-sids** command is used.

- **configure router isis segment-routing prefix-sid-range global**
- **configure router ospf segment-routing prefix-sid-range global**
- **configure router ospf3 segment-routing prefix-sid-range global**

When a shared SID is configured outside the routing instances, it can be used for all instances when the routing protocol is enabled on the interface. The following CLI configures the prefix SIDs.

```
configure
|
+---router
|   +---segment-routing
|       +---sr-mpls
|           +---prefix-sids [<ip-int-name>]
|               no prefix-sids [<ip-int-name>]
|               +---no ipv4-sid
|                   ipv4-sid index <[0..4294967295]>
|                   ipv4-sid label <[32..1048575]>
|               +---no ipv6-sid
|                   ipv6-sid index <[0..4294967295]>
|                   ipv6-sid label <[32..1048575]>
|               ---node-sid
|                   no node-sid
```

The following commands are used for configuration:

- **ipv4-sid**
This command is used to configure the SID associated with the primary IPv4 address of the loopback or system interface.
- **ipv6-sid**
This command is used to configure the SID associated with the primary IPv6 address of the loopback or system interface.
- **node-sid**
This command sets the N-flag. The N-flag is set when the prefix SID is a node SID, as described in RFC 8402. If the N-flag is not set, the address is an SR anycast SID.

The following considerations apply for shared **sr-mpls>prefix-sids**:

- When an **sr-mpls>prefix-sids** is shared between IGP instances, all instances must share the same SR label range. This means that the instances must use the "global" SRGB range.

- Locally configured shared **sr-mpls>prefix-sids** share the statistics on that node, if configured. As a result, when incoming SID statistics on both OSPF and IS-IS are enabled and the SID is shared, the same statistics are displayed for both IGPs.

The following restrictions apply when configuring the **sr-mpls>prefix-sids**:

- The **sr-mpls>prefix-sids** command can only be used for loopback and system interfaces.
- Exporting **sr-mpls>prefix-sids** into BGP and using it for stitching an SR IGP domain with BGP-based SR MPLS tunnels is not supported.
- On the same interface, sharing the node SID across different address families is not allowed (for example, IPv4 node SID in ISIS and IPv6 in OSPFv3 or even IPv4 and IPv6 in the same ISIS instance).
- Configuring a SID as a prefix SID in one instance and as node SID in another instance is not allowed. For example, if IS-IS has assigned an IPv4 node SID or IPv6 node SID to a loopback in an IS-IS instance, OSPF cannot install the same SID on that loopback as a shared **sr-mpls>prefix-sids**.
- Each **sr-mpls/prefix-sids** SID must be unique across all routing instances.
- A regular IGP node SID and SR-MPLS prefix SID can be configured on a single interface for a single IGP algorithm. In this case, the IGP overrides the **configure>router>segment-routing>sr-mpls>prefix-sids** configuration, and only the IGP node SID is advertised.

Use the **show router segment-routing sr-mpls prefix-sids** and **tools dump router segment-routing tunnel** CLI commands to verify the operation of the shared SIDs. For more information, see the *7705 SAR Gen 2 Clear, Monitor, Show, Tools CLI Command Reference Guide*.

Example: Show command output

```
*A:Dut-A# show router segment-routing sr-mpls prefix-sids
```

```
=====
Rtr Base SR-MPLS Prefix-SIDs
=====
```

Interface Name	AF	SID	Label	State
System	IPv4	123	100123	enabled
System	IPv6	234	100234	ifFailed
loopback.0	IPv4	345	100345	ifDown
loopback.0	IPv6	456	100456	ifDown
loopback.4	IPv4	567	100567	failed
loopback.4	IPv6	-	-	adminDown
loopback.6	IPv4	-	-	adminDown
loopback.6	IPv6	678	100678	notPref

```
-----
No. of Prefix-SIDs: 4
=====
```

```
*A:Dut-C# tools dump router segment-routing tunnel
```

```
=====
Legend: (B) - Backup Next-hop for Fast Re-Route
```

```
(D) - Duplicate
```

```
label stack is ordered from top-most to bottom-most
```

```
=====
```

Prefix Sid-Type	Fwd-Type Next Hop(s)	In-Label	Prot-Inst(algoId)	Out-Label(s)	Interface/Tunnel-ID
1.1.1.3 Node	Terminating	20003	IGP-Shared		

```
=====
```

1.1.1.5					
Node	Orig/Transit	20005	ISIS-0		
	10.10.10.2			20005	To_1/1/1(E)
10.10.10.2					
Adjacency	Transit	524287	ISIS-0		
	10.10.10.2			3	To_1/1/1(E)
-----+					
No. of Entries: 3					
-----+					
*A:Dut-C#					

2.1.3 Segment routing shortest path forwarding with IS-IS

This section describes the segment routing shortest path forwarding with IS-IS.

2.1.3.1 IS-IS control protocol changes

The following TLVs and sub-TLVs are defined in *draft-ietf-isis-segment-routing-extensions* and are supported in the implementation of segment routing in IS-IS:

- Prefix Segment Identifier (Prefix-SID) sub-TLV
- Adjacency Segment Identifier (Adj-SID) sub-TLV
- SID/Label Binding TLV
- SR-Capabilities sub-TLV
- SR-Algorithm sub-TLV

This section describes the behaviors of the IS-IS support of the segment routing TLVs and sub-TLVs.

SR OS supports advertising the IS-IS Router Capability TLV (RFC 4971) for topology MT0 and MT2.

Special attention is taken when leaking the IS-IS router capability when both MT0 and MT2 are enabled. When leaking router capability between IS-IS levels, as defined in RFC 7981, a reachability check must be performed. The router performs the following reachability check with MT-IS-IS MT2 enabled:

- leak the router capability TLVs with a valid IPv4 router ID via IPv4 MT0
- for router capabilities with no valid IPv4 router ID but a valid IPv6 router ID, perform a reachability check via MT0 and MT2
- when either an IS-IS IPv4 or IPv6 router ID is reachable, then redistribute router capability is redistributed

If Prefix-SID sub-TLVs for the same prefix are received in different MT numbers of the same IS-IS instance, a tiebreaking mechanism is applied to resolve the Prefix-SID. The IS-IS MT0 and MT2 tiebreaking mechanism sorts a specific prefix and gives precedence as follows:

1. smaller route preference sorts ahead
2. smaller route metric sorts ahead
3. MT0 sorts ahead of MT2
4. if all are equal, the final step is for the smaller IS-IS instance ID to sort ahead

When a duplicate Prefix-SID exists between two different prefixes, an error is logged and a trap is generated, as described in [Error and resource exhaustion handling](#).

The I and V flags are both set to 1 when originating the SR-Capabilities sub-TLV to indicate support for processing both SR MPLS-encapsulated IPv4 and IPv6 packets on the network interfaces of the router. These flags are not checked when the sub-TLV is received. Only the SRGB range is processed.

The algorithm field is set to 0, meaning the Shortest Path First (SPF) algorithm based on the link metric, when originating the SR-Algorithm Capability sub-TLV but is not checked when the sub-TLV is received.

SR OS originates a single Prefix-SID sub-TLV per the IS-IS IP-reachability TLV and processes the first Prefix-SID sub-TLV only if multiple sub-TLVs are received within the same IS-IS IP-reachability TLV.

SR OS encodes the 32-bit index in the Prefix-SID sub-TLV. The 24-bit label is not supported.

Prefix-SID sub-TLV encoding

SR OS originates a Prefix-SID sub-TLV with the following encoding of flags and the following processing rules:

- The R-flag is set if the Prefix-SID sub-TLV, along with its corresponding IP-reachability TLV, is propagated between the levels.
- The N-flag is always set because SR OS supports a Prefix-SID type that is node SID only.
- The P-flag (no-PHP flag) is always set, meaning the label for the Prefix-SID is pushed by the PHP router when forwarding to this router. The SR OS PHP router processes a received Prefix-SID with the P-flag set to 0 and uses implicit-null for the outgoing label toward the router that advertised it, as long as the P-flag is also set to 1.
- The E-flag (Explicit-Null flag) is always set to 0. An SR OS PHP router, however, processes a received Prefix-SID with the E-flag set to 1. When the P-flag is also set to 1, it pushes explicit-null for the outgoing label toward the router that advertised it.
- The V-flag is always set to 0 to indicate an index value for the SID.
- The L-flag is always set to 0 to indicate that the SID index value is not locally significant.
- The algorithm field is always set to 0 to indicate that the SPF algorithm based on the link metric is used and is not checked when the Prefix-SID sub-TLV is received.
- SR OS resolves a Prefix-SID sub-TLV received without the N-flag set but with the prefix length equal to 32. A trap, however, is raised by IS-IS.
- SR OS does not resolve a Prefix-SID sub-TLV received with the N-flag set and a prefix length other than 32. A trap is raised by IS-IS.
- SR OS resolves a Prefix-SID received within an IP-reachability TLV based on the following route preference:
 1. a SID received via Layer 1 in a Prefix-SID sub-TLV part of the IP-reachability TLV
 2. a SID received via Layer 2 in a Prefix-SID sub-TLV part of the IP-reachability TLV
- A prefix received in an IP-reachability TLV is propagated, along with the Prefix-SID sub-TLV, by default from Layer 1 to Layer 2 by an Layer 1/Layer 2 (L1/L2) router. A router in Layer 2 sets up an SR tunnel to the Layer 1 router via the L1/L2 router, which acts as a Label Switching Router (LSR).
- A prefix received in an IP-reachability TLV is not propagated, along with the Prefix-SID sub-TLV, by default from Level 2 to Level 1 by an L1/L2 router. If the user adds a policy to propagate the received prefix, a router in Layer 1 sets up an SR tunnel to the Layer 2 router via the L1/L2 router, which acts as an LSR.

- If a prefix is summarized by an Area Border Router (ABR), the Prefix-SID sub-TLV is not propagated with the summarized route between levels. To propagate the node SID for a /32 prefix, route summarization must be disabled.
- SR OS propagates the Prefix-SID sub-TLV when exporting the prefix to another IS-IS instance; however, it does not propagate if the prefix is exported from a different protocol. When the corresponding prefix is redistributed from another protocol such as OSPF, the prefix SID is removed.

Adj-SID sub-TLV encoding

SR OS originates an Adj-SID sub-TLV with the following encoding of the flags:

- The F-flag is set to 0 to indicate the IPv4 family and is set to 1 for IPv6 family for the adjacency encapsulation.
- The B-flag is set to 0 and is not processed on receipt.
- The V-flag is always set to 1.
- The L-flag is always set to 1.
- The S-flag is set to 0 because assigning an Adj-SID to parallel links between neighbors is not supported. A received Adj-SID with S-flag set is not processed.
- The weight octet is not supported and is set to all zeros.

SID/Label Binding TLV rules and limitations

SR OS can originate the SID/Label Binding TLV as part of the Mapping Server feature (see [Segment routing mapping server function for IPv4 prefixes](#) for more information) for IS-IS MT0 only. Consider the following rules and limitations:

- Only the mapping server Prefix-SID sub-TLV within the TLV is processed and the ILMs installed if the prefixes in the provided range are resolved.
- The range and FEC prefix fields are processed. Each FEC prefix is resolved similar to the Prefix-SID sub-TLV, meaning there must be an IP-reachability TLV received for the exact matching prefix.
- If the same prefix is advertised with both a Prefix-SID sub-TLV and a mapping server Prefix-SID sub-TLV. The resolution follows the following route preference:
 1. SID received via Level 1 in a Prefix-SID sub-TLV part of IP-reachability TLV
 2. SID received via Level 2 in a Prefix-SID sub-TLV part of IP-reachability TLV
 3. SID received via Level 1 in a mapping server Prefix-SID sub-TLV
 4. SID received via Level 2 in a mapping server Prefix-SID sub-TLV
- The entire TLV can be propagated between levels based on the settings of the S-flag. The TLV cannot be propagated between IS-IS instances (see [Segment routing mapping server function for IPv4 prefixes](#) for more information). Finally, a Level 1 or Level 2 router does not propagate the Prefix-SID sub-TLV from the SID/Label Binding TLV (received from a mapping server) into the IP-reachability TLV if the latter is propagated between levels.
- The mapping server that advertised the SID/Label Binding TLV does not need to be in the shortest path for the FEC prefix.
- If the same FEC prefix is advertised in multiple binding TLVs by different routers, the SID in the binding TLV of the first router that is reachable is used. If that router becomes unreachable, the next reachable one is used.

- No check is performed if the content of the binding TLVs from different mapping servers are consistent or not.
- Any other sub-TLV, for example, the SID/Label sub-TLV, ERO metric and unnumbered interface ID ERO, is ignored but the user can view the octets of the received-but-not-supported sub-TLVs using the IGP **show** command.

2.1.3.2 Segment routing multitopology considerations

Segment routing with IS-IS is supported with Multitopology IS-IS MT0 (standard IPv4 and IPv6 topology) and MT2 (IPv6-only topology).

When a user configures the following command, Prefix-SIDs and Adj-SIDs are advertised for MT-ISIS MT2.

```
configure router isis segment-routing multi-topology mt2
```

The encoding the SR-MPLS Prefix-SIDs and Adj-SIDs depends on the combined usage of the following commands.

```
configure router ipv6-routing  
configure router isis multi-topology ipv6-unicast
```

IPv6 advertised in MT0 only

The following logic applies for IPv6 advertised in MT0 only using the following commands:

- **MD-CLI**

```
configure router isis segment-routing multi-topology mt2 false
```

- **classic CLI**

```
configure router isis no segment-routing multi-topology
```

- All protected and unprotected IPv4 and IPv6 Adj-SIDs are advertised in the Traffic Engineering Neighbor (TE-NBR) TLVs.
- For the protected Adj-SIDs a backup is programmed following the MT0 topology, which is the same for both IPv4 and IPv6.

IPv6 advertised in MT2 only

The following logic applies for IPv6 advertised in MT2 only (using the **ipv6-routing mt** and **multi-topology ipv6-unicast** commands):

- IPv4 protected and unprotected Adj-SIDs are advertised in the TE-NBR TLVs.
- IPv6 protected and unprotected Adj-SIDs are advertised in the Multitopology Neighbor (MT-NBR) TLVs.
- For protected Adj-SIDs:
 - For IPv4 protected Adj-SIDs, a backup is programmed following the MT0 topology
 - For IPv6 protected Adj-SIDs, a backup is programmed following the MT2 topology

IPv6 advertised in MT0 and MT2 simultaneously

The following logic applies for IPv6 advertised in MT0 and MT2 simultaneously (using the **ipv6-routing native** and **multi-topology ipv6-unicast** commands):

- IPv4 protected and unprotected Adj-SIDs are advertised in the TE-NBR TLVs.
- IPv6 protected and unprotected Adj-SIDs are advertised in the MT-NBR TLVs.
- IPv6 unprotected Adj-SIDs are also advertised in the TE-NBR TLVs (The advertised unprotected Adj-SID is identical as the IPv6 Adj-SID advertised in MT-NBR TLV of the MT2).
- Programming of backup paths for protected Adj-SIDs:
 - IPv4 Adj-SIDs – a backup is programmed following the MT0 topology
 - IPv6 Adj-SIDs – a backup is programmed following the MT2 topology



Note: No protected IPv6 Adj-SID exists in MT0. Only an unprotected IPv6 Adj-SID exists in MT0.

The **multi-topology mt2** command operates as follows for IPv6 routes:

- **Multi-topology MT2 segment routing disabled**
By default, Segment Routing in multi-topology IS-IS encoding is disabled. Only SR-MPLS tunnels that are IPv6 native SR-MPLS routes (for IPv6 SIDs in MT0) are programmed.
- **Multi-topology MT2 segment routing enabled**
When MT2 segment routing is enabled, the following applies:
 - The **multi-topology mt2** command instructs the IS-IS router to program SR-MPLS tunnels for multi-topology IPv6 routes in MT2.
 - A router configured as IPv6 MT2 only (using the **ipv6-routing mt** and **multi-topology ipv6-unicast** commands), contains only SR-MPLS MT2 topology tunnels that are programmed for IPv6. Only IPv4 tunnels will be programmed for SR-MPLS MT0 topologies.
 - A router configured for both MT0 and IPv6 MT2 (using the **ipv6-routing native** and **multi-topology ipv6-unicast** commands), has IPv6 and IPv4 SR-MPLS tunnels programmed for MT0 and MT2 routes.

The use of IGP shortcuts (RSVP or SR-TE) in MT2 is not supported.

BGP-LS enables the export of multi-topology IS-IS MT2 and MT0 prefixes and Segment Routing SIDs. When TE attributes as defined in RFC 5305 and RFC 8750 are received on a router from remote devices within the MT2 topology, these attributes are seamlessly integrated into the Traffic Engineering Database (TEDB). They are then conveyed through BGP-LS NRLI encoding for dissemination.

IS-IS Link State Packets (LSP) encoding examples

For SR OS, a user can enable Segment Routing MPLS within IS-IS MT0 alone, MT2 alone, or simultaneously in both MT0 and MT2. This choice has a notable effect on how Prefix-SIDs are presented in IS-IS LSPs. The following sections display encoding examples for each of these three scenarios.

IPv6 in MT0

The following applies for IPv6 in MT0.

Traffic Engineering Neighbor (TE-NBR) TLVs advertise all protected and unprotected IPv4 and IPv6 Adj-SIDs.

Example: Encoding for IPv6 in MT0

```

TE IS Nbrs :
  Nbr : Dut-A.00
  Default Metric : 10
  Sub TLV Len : 153
  IF Addr : 1.1.3.3
  IPv6 Addr : 3ffe::101:303
  Nbr IP : 1.1.3.1
  Nbr IPv6 : 3ffe::101:301
  MaxLink BW: 10000000 kbps
  Resvble BW: 10000000 kbps
  Unresvd BW:
    BW[0] : 10000000 kbps
    BW[1] : 10000000 kbps
    BW[2] : 10000000 kbps
    BW[3] : 10000000 kbps
    BW[4] : 10000000 kbps
    BW[5] : 10000000 kbps
    BW[6] : 10000000 kbps
    BW[7] : 10000000 kbps
  Admin Grp : 0x0
  TE Metric : 1000
  TE APP LINK ATTR :
    SABML-flag:Non-Legacy SABM-flags: X
    Delay Min : 1000000 Max : 1000000
    TE Metric : 1000
  Adj-SID: Flags:v4BVL Weight:0 Label:524287
  Adj-SID: Flags:v4VL Weight:0 Label:524285
  Adj-SID: Flags:v6BVL Weight:0 Label:524286
  Adj-SID: Flags:v6VL Weight:0 Label:524284

```

IPv6 in MT2

The following applies for IPv6 in MT2:

- TE-NBR TLVs advertise IPv4 protected and unprotected Adj-SIDs.
- Multitopology Neighbor (MT-NBR) TLVs advertise IPv6 protected and unprotected Adj-SIDs.

Example: Encoding for IPv6 in MT2

```

TE IS Nbrs :
  Nbr : Dut-A.00
  Default Metric : 10
  Sub TLV Len : 103
  IF Addr : 1.1.3.3
  Nbr IP : 1.1.3.1
  MaxLink BW: 10000000 kbps
  Resvble BW: 10000000 kbps
  Unresvd BW:
    BW[0] : 10000000 kbps
    BW[1] : 10000000 kbps
    BW[2] : 10000000 kbps
    BW[3] : 10000000 kbps
    BW[4] : 10000000 kbps
    BW[5] : 10000000 kbps
    BW[6] : 10000000 kbps
    BW[7] : 10000000 kbps
  Admin Grp : 0x0
  TE Metric : 1000
  TE APP LINK ATTR :
    SABML-flag:Non-Legacy SABM-flags: X

```

```

    Delay Min : 1000000 Max : 1000000
    TE Metric : 1000
    Adj-SID: Flags:v4BVL Weight:0 Label:524287
    Adj-SID: Flags:v4VL Weight:0 Label:524281
MT IS Nbrs      :
  MT ID         : 2
  Nbr          : Dut-A.00
  Default Metric : 10
  Sub TLV Len   : 220
  IPv6 Addr     : 3ffe::101:303
  TE APP LINK ATTR :
    SABML-flag:Non-Legacy SABM-flags: X
    Delay Min : 1000000 Max : 1000000
    TE Metric : 1000
  Adj-SID: Flags:v6BVL Weight:0 Label:524286
  Adj-SID: Flags:v6VL Weight:0 Label:524280
  End.X-SID: 300::2000 flags:BP algo:0 weight:0 endpoint:End.X-PSP
  End.X-SID: 310::2000 flags:BP algo:128 weight:0 endpoint:End.X-PSP
  End.X-SID: 320::2000 flags:BP algo:129 weight:0 endpoint:End.X-PSP
  End.X-SID: 330::2000 flags:BP algo:130 weight:0 endpoint:End.X-PSP
  End.X-SID: 340::2000 flags:BP algo:131 weight:0 endpoint:End.X-PSP
  End.X-SID: 350::2000 flags:BP algo:132 weight:0 endpoint:End.X-PSP
  End.X-SID: 360::2000 flags:BP algo:133 weight:0 endpoint:End.X-PSP

```

IPv6 advertised in both MT0 and MT2

The following applies for IPv6 in both MT0 and MT2:

- TE-NBR TLVs advertise IPv4 protected and unprotected Adj-SIDs.
- MT-NBR TLVs advertise IPv6 protected and unprotected Adj-SIDs.
- IPv6 unprotected Adj-SIDs are also advertised in TE-NBR TLVs (in the same Adj-SID).

Example: Encoding for IPv6 advertised in both MT0 and MT2

```

TE IS Nbrs      :
  Nbr          : Dut-A.00
  Default Metric : 10
  Sub TLV Len   : 146
  IF Addr       : 1.1.3.3
  IPv6 Addr     : 3ffe::101:303
  Nbr IP        : 1.1.3.1
  Nbr IPv6      : 3ffe::101:301
  MaxLink BW: 10000000 kbps
  Resvble BW: 10000000 kbps
  Unresvd BW:
    BW[0] : 10000000 kbps
    BW[1] : 10000000 kbps
    BW[2] : 10000000 kbps
    BW[3] : 10000000 kbps
    BW[4] : 10000000 kbps
    BW[5] : 10000000 kbps
    BW[6] : 10000000 kbps
    BW[7] : 10000000 kbps
  Admin Grp : 0x0
  TE Metric : 1000
  TE APP LINK ATTR :
    SABML-flag:Non-Legacy SABM-flags: X
    Delay Min : 1000000 Max : 1000000
    TE Metric : 1000
  Adj-SID: Flags:v4BVL Weight:0 Label:524275
  Adj-SID: Flags:v4VL Weight:0 Label:524273
  Adj-SID: Flags:v6VL Weight:0 Label:524272

```

```

MT IS Nbrs      :
MT ID           : 2
Nbr            : Dut-A.00
Default Metric  : 10
Sub TLV Len     : 220
IPv6 Addr      : 3ffe::101:303
TE APP LINK ATTR :
  SABML-flag:Non-Legacy SABM-flags: X
  Delay Min : 1000000 Max : 1000000
  TE Metric : 1000
Adj-SID: Flags:v6BVL Weight:0 Label:524274
Adj-SID: Flags:v6VL Weight:0 Label:524272
End.X-SID: 300::2000 flags:BP algo:0 weight:0 endpoint:End.X-PSP
End.X-SID: 310::2000 flags:BP algo:128 weight:0 endpoint:End.X-PSP
End.X-SID: 320::2000 flags:BP algo:129 weight:0 endpoint:End.X-PSP
End.X-SID: 330::2000 flags:BP algo:130 weight:0 endpoint:End.X-PSP
End.X-SID: 340::2000 flags:BP algo:131 weight:0 endpoint:End.X-PSP
End.X-SID: 350::2000 flags:BP algo:132 weight:0 endpoint:End.X-PSP
End.X-SID: 360::2000 flags:BP algo:133 weight:0 endpoint:End.X-PSP

```

2.1.3.3 Announcing ELC, MSD-ERLD, and MSD-BMI with IS-IS

IS-IS can announce node the Entropy Label Capability (ELC), Maximum Segment Depth (MSD) for node Entropy Readable Label Depth (ERLD), and the MSD for node Base MPLS Imposition (BMI). If needed, exporting the IS-IS extensions into BGP-LS requires no additional configuration. These extensions are standardized through *draft-ietf-isis-mpls-elc-10*, *Signaling Entropy Label Capability and Entropy Readable Label Depth Using IS-IS*, and RFC 8491, *Signaling Maximum SID Depth (MSD) Using IS-IS*.

When entropy and segment routing are enabled on a router, it automatically announces the ELC, ERLD, and BMI IS-IS values when announcing the IS-IS prefix attributes and router capabilities. The following configuration logic is used:

- The router automatically announces ELC for host prefixes associated with an IPv4 or IPv6 node SID when **segment-routing**, **segment-routing entropy-label**, and **prefix-attributes-tlv** are enabled for IS-IS. Although the ELC capability is a node property, it is assigned to prefixes to allow inter-area or inter-AS signaling. Consequently, the prefix-attribute TLV must be enabled accordingly within IS-IS.
- The router announces the maximum node ERLD for IS-IS when **segment-routing** and **segment-routing entropy-label** are enabled together with **advertise-router-capability**.
- The router announces the maximum node MSD-BMI for IS-IS when **segment-routing** and **advertise-router-capability** are enabled.
- Exporting ELC, MSD-ERLD, and MSD-BMI IS-IS extensions into BGP-LS encoding is enabled automatically when database-export for BGP-LS is configured.
- The announced value for maximum node MSD-ERLD and MSD-BMI can be modified to a smaller number using the **override-bmi** and **override-erld** commands. This can be useful when services (such as EVPN) or more complex link protocols (such as Q-in-Q) are deployed. Provisioning correct ERLD and BMI values helps controllers and local Constrained Shortest Path First (CSPF) to construct valid segment routing label stacks to be deployed in the network.

Use the commands in the following context to configure segment routing parameters.

```
configure router isis segment-routing maximum-sid-depth
```

2.1.3.4 EL for IS-IS segment routing

The router supports the MPLS entropy label (EL), as specified in RFC 6790, on IS-IS segment-routed tunnels. LSR nodes in a network can load balance labeled packets in a more granular way than by hashing on the standard label stack. See the *7705 SAR Gen 2 MPLS Guide* for more information.

The router can announce ELC; however, it cannot process ELC signaling for IS-IS segment-routed tunnels. Instead, ELC is configured at the head-end LER using the **configure router isis entropy-label override-tunnel-elc** command. This command configures the router to ignore any advertisements for ELC that may or may not be received from the network, and instead to assume that the whole domain supports ELs.

2.1.3.5 IPv6 segment routing using MPLS encapsulation

This feature supports SR IPv6 tunnels in IS-IS MT0 and MT2. The user can configure a node SID for the primary IPv6 global address of a loopback interface, which is then advertised in IS-IS. IS-IS automatically assigns and advertises an adjacency SID for each adjacency with an IPv6 neighbor. After the node SID is resolved, it is used to install an IPv6 SR-ISIS tunnel in the TTM for use by the services.

2.1.3.5.1 IS-IS MT0 and MT2 extensions

The IS-IS extensions support the advertising and resolution of the prefix SID sub-TLV within the IP reach TLV-236 (IPv6 MT0) or the IP Reach TLV-237 (MT2), as defined in RFC 8667, *IS-IS Extensions for Segment Routing*. The adjacency SID is still advertised as a sub-TLV of the Extended IS Reachability TLV 22 (MT0) or TLV 222 (MT2), as defined in RFC 8667. The router sets the V-flag and I-flag in the SR-capabilities sub-TLV to indicate that it can process SR MPLS-encapsulated IPv4 and IPv6 packets on its network interfaces.

For more details of the processing of the prefix SID and adjacency SID sub-TLVs, see [IS-IS control protocol changes](#).

2.1.3.5.2 Service and routing contexts supported

IPv6 SR tunnels support the same services and routing contexts as in IPv4 SR tunnel. For more details see [BGP shortcuts using segment routing tunnels](#), [BGP labeled route resolution using segment routing tunnels](#), and [Service packet forwarding with segment routing](#).

2.1.3.5.3 Services using SDP with an SR IPv6 tunnel

The MPLS SDP of type **sr-isis** with a **far-end** option using an IPv6 address is supported. Note the SDP must have the same IPv6 **far-end** address, used by the control plane for the T-LDP session, as the prefix of the node SID of the SR IPv6 tunnel.

```
configure
- service
  - [no] sdp sdp-id mpls
  - [no] far-end ipv6-address
  - sr-isis
  - no sr-isis
```

The **bgp-tunnel**, **lsp**, **sr-te lsp**, **sr-ospf**, and **mixed-lsp-mode** commands are blocked within the SDP configuration context when the far end is an IPv6 address.

SDP admin groups are not supported with an SDP using an SR IPv6 tunnel, or with SR-OSPF for IPv6 tunnels, and the attempt to assign them is blocked in the CLI.

Services that use LDP control plane such as T-LDP VPLS and R-VPLS, VLL, and IES/VPRN spoke interface have the spoke SDP (PW) signaled with an IPv6 T-LDP session because the **far-end** option is configured to an IPv6 address. The spoke SDP for these services binds to an SDP that uses an SR IPv6 tunnel where the prefix matches the **far-end** address. SR OS also supports the following:

- the IPv6 PW control word with both data plane packets and VCCV OAM packets
- hash label and entropy label, with the above services
- network domains in VPLS

The PW switching feature is not supported with LDP IPv6 control planes. As a result, the CLI does not allow the user to enable the **vc-switching** option whenever one or both spoke SDPs uses an SDP that has the **far-end** configured as an IPv6 address.

L2 services that use BGP control plane such as dynamic MS-PW, BGP-AD VPLS, BGP-VPLS, BGP-VPWS, and EVPN MPLS cannot bind to an SR IPv6 tunnel because a BGP session to a BGP IPv6 peer does not support advertising an IPv6 next hop for the L2 NLRI. As a result, these services do not auto-generate SDPs using an SR IPv6 tunnel. In addition, they skip any provisioned SDPs with **far-end** configured to an IPv6 address when the **use-provisioned-sdp** option is enabled.

SR OS also supports multi homing with T-LDP active/standby FEC 128 spoke SDP using SR IPv6 tunnel to a VPLS/B-VPLS instance. BGP multi homing is not supported because BGP IPv6 does not support signaling an IPv6 next hop for the L2 NLRI.

The Shortest Path Bridging (SPB) feature works with spoke SDPs bound to an SDP that uses an SR IPv6 tunnel.

2.1.3.6 Segment routing mapping server function for IPv4 prefixes

The mapping server feature supports the configuration and advertisement, in IS-IS, of the node SID index for prefixes of routers in the LDP domain. This is performed in the router acting as a mapping server and using a prefix-SID sub-TLV within the SID/label binding TLV in IS-IS.

Use the following command syntax to configure the SR mapping database in IS-IS:

```
configure
  - router
    - [no] isis
      - segment-routing
      - no segment-routing
      - mapping-server
        - sid-map node-sid {index 0..4294967295 [range 0..65535]} prefix {{ip-
address/mask} | {ip-address}{netmask}} [set-flags {s}] [level {1 | 2 | 1/2}]
        - no sid-map node-sid index 0..4294967295
```

The user enters the node SID index, for one prefix or a range of prefixes, by specifying the first index value and, optionally, a range value. The default value for the range option is 1. Only the first prefix in a consecutive range of prefixes must be entered. If the user enters the first prefix with a mask lower than 32, the SID/label binding TLV is advertised, but a router that receives it does not resolve the prefix SID and instead generates a trap.

The **no** form of the **sid-map** command deletes the range of node SIDs beginning with the specified index value. The **no** form of the **mapping-server** command deletes all node SID entries in the IS-IS instance.

The S-flag indicates to the IS-IS routers in the network that the flooding scope of the SID/label binding TLV is the entire domain. In that case, a router receiving the TLV advertisement leaks it between IS-IS levels. If leaked from Level 2 to Level 1, the D-flag must be set; this prevents the TLV from being leaked back into level 2. Otherwise, the S-flag is clear by default and routers receiving the mapping server advertisement do not leak the TLV.



Note: SR OS does not leak this TLV between IS-IS instances and does not support the multitopology SID/label binding TLV format. In addition, the user can specify the flooding scope of the mapping server for the generated SID/label binding TLV using the **level** option. This option allows further narrowing of the flooding scope configured under the router IS-IS level-capability for one or more SID/label binding TLVs if required. The default flooding scope of the mapping server is L1 or L2, which can be narrowed by what is configured under the router IS-IS level-capability.

The A-flag indicates that a prefix for which the mapping server prefix SID is advertised is directly attached. The M-flag advertises a SID for a mirroring context to provide protection against the failure of a service node. None of these flags are supported on the mapping server; the mapping client ignores them.

Each time a prefix or a range of prefixes is configured in the SR mapping database in any routing instance, the router issues for this prefix, or range of prefixes, a prefix-SID sub-TLV within an IS-IS SID/label binding TLV in that instance. The flooding scope of the TLV from the mapping server is determined as previously described. No further check of the reachability of that prefix in the mapping server route table is performed. No check of the SID index is performed to determine whether the SID index is a duplicate of an existing prefix in the local IGP instance database or if the SID index is out of range with the local SRGB.

2.1.3.6.1 IP prefix resolution for segment routing mapping server

The following processing rules apply for IP prefix resolution:

- SPF calculates the next hops, up to **max-ecmp**, to reach a destination node.
- Each prefix inherits the next hops of one or more destination nodes advertising it.
- A prefix advertised by multiple nodes, all reachable with the same cost, inherits up to **max-ecmp** next hops from the advertising nodes.
- The next-hop selection value, up to **max-ecmp**, is based on sorting the next hops by:
 - lowest next-hop router ID
 - lowest interface index, for parallel links to same router ID

Each next hop keeps a reference to the destination nodes from which it was inherited.

2.1.3.6.2 Prefix SID resolution for segment routing mapping server

This section describes the processing rules for prefix SID resolution.

- For a specific prefix, IGP selects the SID value among multiple advertised values in the following order:
 1. the local intra-area SID owned by this router
 2. the prefix SID sub-TLV advertised within an IP reach TLV

If multiple SIDs exist, the IGP selects the SID corresponding to the destination router or the ABR with the lowest system ID that is reachable using the first next hop of the prefix.

3. the IS-IS SID and label binding TLV from the mapping server

If multiple SIDs exist, the IGP selects the following, using the preference rules in *draft-ietf-spring-conflict-resolution-05* when applied to the SRMS entries of the conflicting SIDs. The order of these rules is as follows:

- a. smallest range
- b. smallest starting address
- c. smallest algorithm
- d. smallest starting SID



Note: If an L1L2 router acts as a mapping server and also re-advertises the mapping server prefix SID from other mapping servers, the redistributed mapping server prefix SID is preferred by other routers resolving the prefix, which may result in not selecting the mapping server respecting these rules.

- The selected SID is used with all ECMP next hops from the IP prefix resolution in step toward all destination nodes or ABR nodes that advertised the prefix.
- If duplicate prefix SIDs exist for different prefixes after these processing steps are completed, the first SID that is processed is programmed according to its corresponding prefix. Subsequent SIDs cause a duplicate SID trap message and are not programmed. The corresponding prefixes are still resolved and programmed normally using IP next-next-hops.

2.1.3.6.3 SR tunnel programming for segment routing mapping server

The following processing rules apply for SR tunnel programming:

- If the prefix SID is resolved from a prefix SID sub-TLV advertised within an IP Reachability TLV, one of the following applies:
 - The SR ILM label is swapped to an SR NHLFE label, as in SR tunnel resolution when the next hop of the IS-IS prefix is SR-enabled.
 - The SR ILM label is stitched to an LDP FEC of the same prefix when either the next hop of the IS-IS prefix is not SR-enabled (no SR NHLFE) or an import policy rejects the prefix (SR NHLFE is deprogrammed).
The LDP FEC can also be resolved by using the same or a different IGP instance as that of the prefix SID sub-TLV or by using a static route.
- If the prefix SID is resolved from a mapping server advertisement, one of the following applies:
 - The SR ILM label is stitched to an LDP FEC of the same prefix, if one exists. The stitching is performed even if an import policy rejects the prefix in the local IS-IS instance.
The LDP FEC can also be resolved by using a static route, a route within an IS-IS instance, or a route within an OSPF instance. The IS-IS or OSPF instances can be the same as, or different from the IGP instance that advertised the mapping server prefix SID sub-TLV.
 - The SR ILM label is swapped to an SR NHLFE label. This is only possible if a route is exported from another IGP instance into the local IGP instance without propagating the prefix SID sub-TLV with the route. Otherwise, the SR ILM label is swapped to an SR NHLFE label toward the stitching node.

2.1.4 Segment routing shortest path forwarding with OSPF

This section describes the segment routing shortest path forwarding with OSPF.

2.1.4.1 OSPFv2 control protocol changes

The following TLVs and sub-TLVs are defined in *draft-ietf-ospf-segment-routing-extensions-04* and are required for the implementation of segment routing in OSPF:

- the prefix SID sub-TLV part of the OSPFv2 Extended Prefix TLV
- the prefix SID sub-TLV part of the OSPFv2 Extended Prefix Range TLV
- the adjacency SID sub-TLV part of the OSPFv2 Extended Link TLV
- SID/Label Range capability TLV
- SR-Algorithm capability TLV

This section describes the behaviors and limitations of OSPF support of segment routing TLVs and sub-TLVs.

SR OS originates a single prefix SID sub-TLV per OSPFv2 Extended Prefix TLV and processes the first one only if multiple prefix SID sub-TLVs are received within the same OSPFv2 Extended Prefix TLV.

SR OS encodes the 32-bit index in the prefix SID sub-TLV. The 24-bit label or variable IPv6 SID is not supported.

SR OS originates a prefix SID sub-TLV with the following encoding of the flags:

- The NP-Flag is always set. The label for the prefix SID is pushed by the PHP router when forwarding to this router. The SR OS PHP router processes a received prefix SID with the NP-flag set to zero and uses implicit-null for the outgoing label toward the router that advertised it.
- The M-Flag is always unset because SR OS does not support originating a mapping server prefix-SID sub-TLV.
- The E-flag is always set to zero. An SR OS PHP router, however, processes a received prefix SID with the E-flag set to 1, and when the NP-flag is also set to 1, it pushes explicit-null for the outgoing label toward the router that advertised it.
- The V-flag is always set to 0 to indicate an index value for the SID.
- The L-flag is always set to 0 to indicate that the SID index value is not locally significant.
- The algorithm field is set to zero to indicate Shortest Path First (SPF) algorithm based on link IGP metric or to the flexible algorithm number.

SR OS resolves a prefix SID received within an Extended Prefix TLV based on the following route preference:

- SID received via an intra-area route in a prefix SID sub-TLV part of the Extended Prefix TLV
- SID received via an inter-area route in a prefix SID sub-TLV part of the Extended Prefix TLV

SR OS originates an adjacency SID sub-TLV with the following encoding of the flags:

- The B-flag is set to zero and is not processed on receipt.
- The V-flag is always set.
- The L-flag is always set.

- The G-flag is not supported.
- The weight octet is not supported and is set to all zeros.

An adjacency SID is assigned to next hops over both the primary and secondary interfaces.

SR OS can originate the OSPFv2 Extended Prefix Range TLV as part of the Mapping Server feature and can process it properly, if received. Consider the following rules and limitations:

- Only the prefix SID sub-TLV within the TLV is processed and the ILMs are installed if the prefixes are resolved.
- The range and address prefix fields are processed. Each prefix is resolved separately.
- If the same prefix is advertised with both a prefix SID sub-TLV in an IP-reachability TLV and a mapping server Prefix-SID sub-TLV, the resolution follows the following route preference:
 - the SID received via an intra-area route in a prefix SID sub-TLV part of Extended Prefix TLV
 - the SID received via an inter-area route in a prefix SID sub-TLV part of Extended Prefix TLV
 - the SID received via an intra-area route in a prefix SID sub-TLV part of a OSPFv2 Extended Range Prefix TLV
 - the SID received via an inter-area route in a prefix SID sub-TLV part of a OSPFv2 Extended Range Prefix TLV
- No leaking of the entire TLV is performed between areas. An ABR does not propagate the prefix-SID sub-TLV from the Extended Prefix Range TLV into an Extended Prefix TLV if the latter is propagated between areas.
- The mapping server which advertised the OSPFv2 Extended Prefix Range TLV does not need to be in the shortest path for the FEC prefix.
- If the same FEC prefix is advertised in multiple OSPFv2 Extended Prefix Range TLVs by different routers, the SID in the TLV on the first router that is reachable is used. If that router becomes unreachable, the next reachable one is used.
- There is no check to determine whether the contents of the OSPFv2 Extended Prefix Range TLVs received from different mapping servers are consistent.
- Any other sub-TLV (for example, the ERO metric and unnumbered interface ID ERO) is ignored, but the user can use the IGP **show** command to see the octets of the received but not supported sub-TLVs.

SR OS supports propagation on the ABR of external prefix LSAs into other areas with routeType set to 3 as per *draft-ietf-ospf-segment-routing-extensions-04*.

SR OS supports propagation on the ABR of external prefix LSAs with route type 7 from a Not-So-Stubby Area (NSSA) into other areas with route type set to 5 as per *draft-ietf-ospf-segment-routing-extensions-04*. SR OS does not support propagating the prefix SID sub-TLV between OSPF instances.

When the user configures an OSPF import policy, the outcome of the policy applies to prefixes resolved in the RTM and the corresponding tunnels in the TTM. A prefix removed by the policy does not appear as both a route in the RTM and as a segment routing tunnel in the TTM.

2.1.4.2 OSPFv3 control protocol changes

The OSPFv3 extensions support the following TLVs:

- **a prefix SID that is a sub-TLV of the OSPFv3 prefix TLV**

The OSPFv3 prefix TLV is a new top-level TLV of the extended prefix LSA introduced in *draft-ietf-ospf-ospfv3-lsa-extend*. The OSPFv3 instance can operate in either LSA sparse mode or extended LSA mode.

The **config>router>extended-lsa only** command advertises the prefix SID sub-TLV in the extended LSA format in both cases.

- **an adjacency SID that is a sub-TLV of the OSPFv3 router-link TLV**

The OSPFv3 router-link TLV is a new top-level TLV in the extended router LSA introduced in *draft-ietf-ospf-ospfv3-lsa-extend*. The OSPFv3 instance can operate in either LSA sparse mode or extended LSA mode. The **config>router>extended-lsa only** command advertises the adjacency SID sub-TLV in the extended LSA format in both cases.

- **the SR-Algorithm TLV and the SID/Label range TLV**

Both of these TLVs are part of the TLV-based OSPFv3 Router Information Opaque LSA defined in RFC 7770.

2.1.4.3 Announcing ELC, MSD-ERLD, and MSD-BMI with OSPF

OSPF can announce node ELC, MSD for node ERLD, and the MSD for node BMI. If needed, exporting these OSPF extensions into BGP-LS requires no additional configuration. These extensions are standardized through *draft-ietf-ospf-mpls-elc-12*, *Signaling Entropy Label Capability and Entropy Readable Label-stack Depth Using OSPF*, and RFC 8476, *Signaling Maximum SID Depth (MSD) Using OSPF*.

When entropy and segment routing are enabled on a router, it automatically announces the ELC, ERLD, and BMI OSPF values. The following configuration logic is used:

- The router automatically announces ELC for host prefixes associated with a node SID when **segment-routing** and **segment-routing entropy-label** are enabled for OSPF.
- The router announces the maximum node ERLD for OSPF when **segment-routing** and **segment-routing entropy-label** are enabled together with **advertise-router-capability**.
- The router announces the maximum node MSD-BMI for OSPF when **segment-routing advertise-router-capability** are enabled.
- Exporting ELC, MSD-ERLD and MSD-BMI OSPF extensions into BGP-LS encoding occurs automatically when **database-export** for BGP-LS is configured.
- The announced value for maximum node MSD-ERLD and MSD-BMI can be modified to a smaller number using the **override-bmi** and **override-erld** commands. This can be useful when services (such as EVPN) or more complex link protocols (such as Q-in-Q) are deployed. Provisioning correct ERLD and BMI values helps controllers and local CSPF to construct valid segment routing label stacks to be deployed in the network.

Use the commands in the following context to configure segment routing parameters.

```
configure router ospf segment-routing maximum-sid-depth
```

2.1.4.4 EL for OSPF segment routing

The router supports the MPLS EL, as specified in RFC 6790, on OSPF segment-routed tunnels. LSR nodes in a network can load-balance labeled packets in a more granular way than by hashing on the standard label stack. See the *7705 SAR Gen 2 MPLS Guide* for more information.

The router can announce ELC; however, it cannot process ELC signaling for OSPF segment-routed tunnels. Instead, ELC is configured at the head-end LER using the **configure router ospf entropy-label override-tunnel-elc** command. This command configures the router to ignore any advertisements for ELC that may or may not be received from the network, and to assume that the whole domain supports ELs.

2.1.4.5 Segment routing mapping server for IPv4 prefixes

The mapping server feature configures and advertises, in OSPF, the node SID index for prefixes of routers in the LDP domain. This is performed in the router acting as a mapping server and using a prefix-SID sub-TLV within an OSPF Extended Prefix Range TLV.

Use the following command syntax to configure the SR mapping database in OSPF:

```
configure
- router
  - [no] ospf
    - segment-routing
    - no segment-routing
    - mapping-server
      - sid-map node-sid {index 0 to 4294967295 [range 1 to 65535]} prefix
        {{ip-address/mask}}[netmask]][scope {area area-id | as}]
      - no sid-map node-sid index 0 to 4294967295
```

The user enters the node SID index, for one prefix or a range of prefixes, by specifying the first index value and, optionally, a range value. The default value for the **range** option is 1. Only the first prefix in a consecutive range of prefixes must be entered. If the user enters the first prefix with a mask lower than 32, the OSPF Extended Prefix Range TLV is advertised, but a router that receives the OSPF Extended Prefix Range TLV does not resolve the SID and instead generates a trap.

The **no** form of the **sid-map** command deletes the range of node SIDs beginning with the specified index value. The **no** form of the **mapping-server** command deletes all node SID entries in the OSPF instance.

Use the **scope** option to specify the flooding scope of the mapping server for the generated OSPF Extended Prefix Range TLV. There is no default value. If the scope is a specific area, the TLV is flooded only in that area.

An ABR that propagates an intra-area OSPF Extended Prefix Range TLV flooded by the mapping server in that area into other areas sets the inter-area flag (IA-flag). The ABR also propagates the TLV if it is received with the IA-flag set from other ABR nodes but only from the backbone to leaf areas and not leaf areas to the backbone. However, if the identical TLV was advertised as an intra-area TLV in a leaf area, the ABR does not flood the inter-area TLV into that leaf area.



Note: SR OS does not leak the OSPF Extended Prefix Range TLV between OSPF instances.

Each time a prefix or a range of prefixes is configured in the SR mapping database in any routing instance, the router issues for this prefix, or range of prefixes, a prefix-SID sub-TLV within an OSPF Extended Prefix Range TLV in that instance. The flooding scope of the TLV from the mapping server is determined as previously described. The reachability of that prefix in the mapping server route table is not checked. Additionally, the SR OS does not check whether the SID index is a duplicate of an existing prefix in the local IGP instance database or if the SID index is out of range with the local SRGB.

2.1.4.5.1 IP prefix resolution for segment routing mapping server

The following processing rules apply for IP prefix resolution:

- SPF calculates the next hops, up to **max-ecmp**, to reach a destination node.
- Each prefix inherits the next hops of one or more destination nodes advertising it.
- A prefix advertised by multiple nodes, all reachable with the same cost, inherits up to **max-ecmp** next hops from the advertising nodes.
- The next-hop selection value, up to **max-ecmp**, is based on sorting the next hops by:
 - lowest next-hop router ID
 - lowest interface index, for parallel links to same router ID

Each next hop keeps a reference to the destination nodes from which it was inherited.

2.1.4.5.2 Prefix SID resolution for segment routing mapping server

The following processing rules apply for prefix SID resolution:

- For a specific prefix, IGP selects the SID value among multiple advertised values in the following order:
 1. local intra-area SID owned by this router
 2. prefix SID sub-TLV advertised within a OSPF Extended Prefix TLV

If multiple SIDs exist, select the SID corresponding to the destination router or ABR with the lowest OSPF router ID which is reachable via the first next hop of the prefix
 3. OSPF Extended Prefix Range TLV from mapping server

If multiple SIDs exist, select the following, using the preference rules in *draft-ietf-spring-conflict-resolution-05* when applied to the SRMS entries of the conflicting SIDs. The order of these rules is as follows:

 - a. smallest range
 - b. smallest starting address
 - c. smallest algorithm
 - d. smallest starting SID
- The selected SID is used with all ECMP next hops from step 1 toward all destination nodes or ABR nodes which advertised the prefix.
- If duplicate prefix SIDs exist for different prefixes after above steps, the first SID which is processed is programmed for its corresponding prefix. Subsequent SIDs causes a duplicate SID trap message and are not programmed. The corresponding prefixes are still resolved normally using IP next hops.

2.1.4.5.3 SR tunnel programming for segment routing mapping server

The following processing rules apply for SR tunnel programming:

- If the prefix SID is resolved from a prefix SID sub-TLV advertised within an OSPF Extended Prefix TLV, one of the following applies.

- The SR ILM label is swapped to an SR NHLFE label as in SR tunnel resolution when the next hop of the OSPF prefix is SR-enabled.
- The SR ILM label is stitched to an LDP FEC of the same prefix when either the next hop of the OSPF prefix is not SR enabled (no SR NHLFE) or an import policy rejects the prefix (SR NHLFE deprogrammed).

The LDP FEC can also be resolved using the same or a different IGP instance as that of the prefix SID sub-TLV or using a static route.

- If the prefix SID is resolved from a mapping server advertisement, one of the following applies.

- The SR ILM label is stitched to an LDP FEC of the same prefix, if one exists. The stitching is performed even if an import policy rejects the prefix in the local OSPF instance.

The LDP FEC can also be resolved using a static route, a route within an OSPF instance, or a route within an OSPF instance. The latter two can be the same as, or different from the IGP instance that advertised the mapping server prefix SID sub-TLV.

- The SR ILM label is swapped to an SR NHLFE label toward the stitching node.

2.1.5 Segment routing with BGP

Segment routing allows a router, potentially by action of an SDN controller, to source route a packet by prepending a segment router header containing an ordered list of SIDs. Each SID can be viewed as a topological or service-based instruction. A SID can have a local impact to one particular node or it can have a global impact within the SR domain, such as the instruction to forward the packet on the ECMP-aware shortest path to reach a prefix, "P". With SR-MPLS, each SID is an MPLS label and the complete SID list is a stack of labels in the MPLS header.

To ensure that all the routers in a network domain have a common understanding of a topology SID, the association of the SID with an IP prefix must be propagated by a routing protocol. Traditionally, this is done by an IGP protocol; however, in some cases, the meaning of a SID may need to be propagated across network boundaries that extend beyond IGP protocol boundaries. For these cases, BGP can carry the association of an SR-MPLS SID with an IP prefix by attaching a prefix SID BGP path attribute to an IP route belonging to a labeled-unicast address family.

The prefix SID attribute attached to a labeled-unicast route for prefix P advertises a SID corresponding to the network-wide instruction to forward the packet along the ECMP-aware BGP-computed best path or paths to reach P. The prefix SID attribute is an optional transitive BGP path attribute with type code 40. This attribute encodes a 32-bit label index into the SRGB space and can provide details about the SRGB space of the originating router. The encoding of this BGP path attribute and its semantics are further described in *draft-ietf-idr-bgp-prefix-sid*.

Using the **block-prefix-sid** BGP command, an SR OS router with upgraded software that processes the prefix SID attribute can prevent it from propagating outside the segment routing domain where it is applicable. The **block-prefix-sid** command removes the prefix SID attribute from all routes sent and received to and from the iBGP and eBGP peers included in the scope of the command. By default, the attribute propagates without restriction.

SR OS attaches a meaning to a prefix SID attribute only when it is attached to routes belonging to the labeled-unicast IPv4 and labeled-unicast IPv6 address families. When attached to routes of unsupported address families, the prefix SID attribute is ignored but still propagated, as with any other optional transitive attribute.

Segment routing must be administratively enabled under BGP using the **config router bgp segment-routing no shutdown** command. When segment routing is configured, the following considerations apply:

- For BGP to redistribute a static or IGP route for a /32 IPv4 prefix as a label-ipv4 route, or a /128 IPv6 prefix as a label-ipv6 route, with a prefix SID attribute, a **route-table-import** policy with an **sr-label-index** action is required.
- For BGP to add or modify the prefix SID attribute in a received label-ipv4 or label-ipv6 route, a BGP **import** policy with an **sr-label-index** action is required.
- For BGP to advertise a label-ipv4 or label-ipv6 route with an incoming datapath label based on the attached prefix SID attribute when BGP segment routing is disabled, new label values assigned to label-ipv4 or label-ipv6 routes come from the dynamic label range of the router and have no network-wide impact.

To enable BGP segment routing, the Base router BGP instance must be associated with a **prefix-sid-range**. This command specifies which SRGB label block to use (for example, to allocate labels). This command also specifies which SRGB label block to advertise in the Originator SRGB TLV of the prefix SID attribute. The **global** parameter value indicates that BGP should use the SRGB as configured under **config>router>mpls-labels>sr-labels**. The **start-label** and **max-index** parameters are used to restrict the BGP prefix SID label range to a subset of the global SRGB.



Note: The **start-label** and **max-index** values must be within the global SRGB range or the command fails.

This is useful when partitioning of the SRGB into non-overlapping subranges dedicated to different IGP/BGP protocol instances is required. Segment routing under BGP must be shutdown before any changes can be made to the **prefix-sid-range** command.

A unique label-index value is assigned to each unique IPv4 or IPv6 prefix that is advertised with a BGP prefix SID. If label-index N1 is assigned to a BGP-advertised prefix P1, and N1 plus the SRGB start label creates a label value that conflicts with another SR programmed LFIB entry, the conflict situation is addressed according to the following rules:

- If the conflict is with another BGP route for prefix P2 that was advertised with a prefix SID attribute, all the conflicting BGP routes for P1 and P2 are advertised with a normal BGP-LU label from the dynamic label range.
- If the conflict is with an IGP route and BGP is not attempting to redistribute that IGP route as a label-ipv4 or label-ipv6 route with a route-table-import policy action that uses the **prefer-igp** keyword in the **sr-label-index** command, the IGP route takes priority and the BGP route is advertised with a normal BGP-LU label from the dynamic label range.
- If the conflict is with an IGP route and BGP is attempting to redistribute that IGP route as a label-ipv4 or label-ipv6 route with a route-table-import policy action that uses the **prefer-igp** keyword in the **sr-label-index** command, this is not considered a conflict and BGP uses the IGP-signaled label-index to derive its advertised label. This has the effect of stitching the BGP segment routing tunnel to the IGP segment routing tunnel.



Note: This use of the **prefer-igp** option is only possible when BGP segment routing is configured with the **prefix-sid-range global** command.

Any /32 label-ipv4 or /128 label-ipv6 BGP routes containing a prefix SID attribute are resolvable and used in the same way as /32 label-ipv4 or /128 label-ipv6 routes without a prefix SID attribute. These routes are installed in the route table and tunnel table (unless **disable-route-table-install** or **selective-label-ipv4-install** are enabled). These routes can have ECMP next hops or FRR backup next hops and be used as transport tunnels for any service that supports BGP-LU transport.



Note: Receiving a /32 label-ipv4 or /128 label-ipv6 route with a prefix SID attribute does not create a tunnel in the segment-routing database; it only creates a label swap entry when the route is re-advertised with a new next hop.

It is recommended the first SID in any SID-list of an SR policy should not be based on a BGP prefix SID; if this recommendation is not followed, then the SID-list may appear to be valid but the datapath is not programmed correctly. However, it is acceptable to use a BGP prefix SID for any SID other than first SID in any SR policy.

2.1.6 Segment routing operational procedures

This section describes the segment routing operational procedures.

2.1.6.1 Prefix advertisement and resolution

After segment routing is successfully enabled in the IS-IS or OSPF instance, the router performs the following operations:

1. The router advertises the Segment Routing Capability sub-TLV to routers in all areas or levels of this IGP instance. Only neighbors with which the router established an adjacency can interpret the SID and label range information and use it for calculating the label to swap to or push for a specific resolved prefix SID.
2. The router advertises the assigned index for each configured node SID in the new prefix SID sub-TLV with the N-flag (node SID flag) set. The segment routing module then programs the ILM with a pop operation for each local node SID in the datapath.
3. The router assigns and advertise an adjacency SID label for each formed adjacency over a network IP interface in the Adjacency SID sub-TLV, according to the following rules and limitations:
 - The Adjacency SID sub-TLV is advertised for both numbered and unnumbered network IP interfaces.
 - The Adjacency SID is not advertised for an IES interface because access interfaces do not support MPLS.
 - The Adjacency SID sub-TLV must be unique per instance and per adjacency.

IS-IS can establish an adjacency for both IPv4 (MT0) and IPv6 (MT0 or MT2) address families over the same link. In this case, a different adjacency SID is assigned to each next hop. However, the existing IS-IS implementation assigns a single Protect-Group ID (PG-ID) to the adjacency and therefore when the state machine of a BFD session tracking the IPv4 or IPv6 next hop times out, an action is triggered for the prefixes of both address families over that adjacency.

The segment routing module programs the ILM with a swap to an implicit null label operation for each advertised adjacency SID.

4. The router resolves received prefixes. If a prefix SID sub-TLV exists, the segment routing module programs the ILM with a swap operation and an LTN with a push operation, both pointing to the primary/LFA NHLFE. A segment routing tunnel is also added to the TTM. If a node SID resolves over an IES interface, the datapath is not programmed and a trap message is generated. Only next-hops of an ECMP set corresponding to network IP interfaces are programmed in the datapath; next-hops corresponding to IES interfaces are not programmed. If the user configures the interface as network on one side and IES on the other side, MPLS packets for the segment routing tunnel received on the access side are dropped.



Note: LSA filtering causes SIDs not to be sent in one direction, which means that some node SIDs are resolved in parts of the network upstream of the advertisement suppression.

When the user enables segment routing in an IGP instance, the main SPF and LFA SPF are computed normally and the primary next-hop and LFA backup next-hop for a received prefix are added to RTM without the label information advertised in the prefix SID sub-TLV. In all cases, the SR tunnel is not added into RTM.

See the following sections for more information about all TLVs and sub-TLVs for both IS-IS and OSPF protocols.

- [IS-IS control protocol changes](#)
- [OSPFv2 control protocol changes](#)
- [OSPFv3 control protocol changes](#)

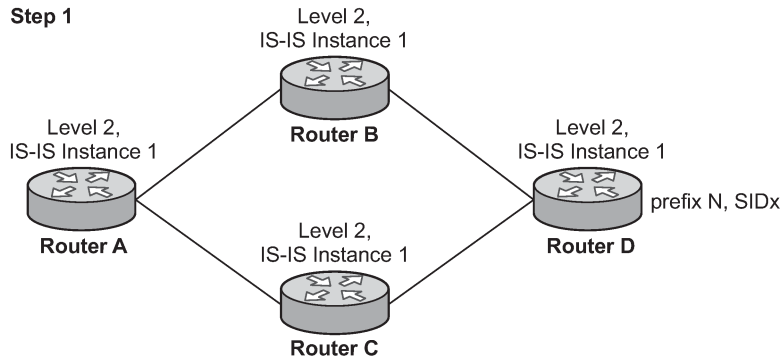
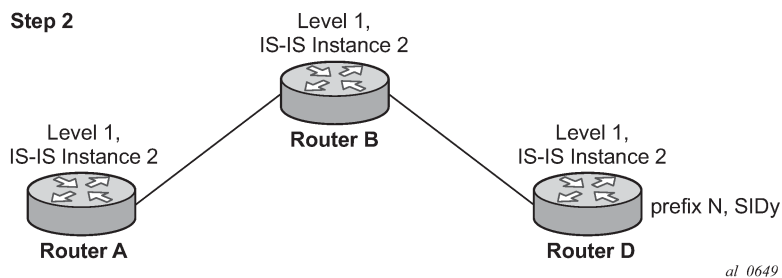
2.1.6.2 Error and resource exhaustion handling

The router performs the procedures described in the following sections when resolving a node SID prefix.

2.1.6.2.1 Supporting multiple topologies for the same destination prefix

SR OS can assign different prefix-SID indexes and labels to the same prefix in different IGP instances. While other routers that receive these prefix SIDs program a single route into the RTM based on the winning instance ID as per RTM route type preference, SR OS adds two tunnels to this destination prefix in the TTM. This supports multiple topologies for the same destination prefix. [Figure 2: Programming multiple tunnels to the same destination](#) shows two different instances (Level 2, IS-IS instance 1 and Level 1, IS-IS instance 2), where Router D has the same prefix destination with different SIDs (SIDx and SIDy).

Figure 2: Programming multiple tunnels to the same destination

Step 1**Step 2**

al_0649

Assume the following route-type preference in the RTM and tunnel-type preference in the TTM are configured:

- ROUTE_PREF_ISIS_L1_INTER (RTM) 15
- ROUTE_PREF_ISIS_L2_INTER (RTM) 18
- ROUTE_PREF_ISIS_TTM 10



Note: The TTM tunnel type preference is not used by the segment routing module. It is put in the TTM and is used by other applications, such as VPRN auto-bind and BGP shortcut, to select a TTM tunnel.

1. Router A performs the following resolution within the single Level 2, IS-IS instance 1. All metrics are the same and ECMP = 2.
 - For prefix N, the RTM entry is:
 - prefix N
 - nhop1 = B
 - nhop2 = C
 - preference 18
 - For prefix N, the SR tunnel TTM entry is:
 - tunnel-id 1: prefix N-SIDx
 - nhop1 = B
 - nhop2 = C
 - tunl-pref 10

2. Add Level 1, IS-IS instance 2 in the same configuration, but in routers A, B, and C only.

- For prefix N, the RTM entry is:

- prefix N
- nhop1 = B
- preference 15

The RTM prefers Level 1 route over Level 2 route.

- For prefix N, there are two SR tunnel entries in TTM:

SR entry for Level 2:

- tunnel-id 1: prefix N-SIDx
- nhop1 = B
- nhop2= C
- tunl-pref 10

The SR entry for Level 1 is tunnel-id 2: prefix N-SIDy.

2.1.6.2.2 Resolving received SID indexes or labels to different routes of the same prefix within the same IGP instance

The router can perform the following variations of this procedure:

- When the SR OS does not allow assigning the same SID index or label to different routes of the same prefix within the same IGP instance, the router resolves only one of the duplicate SIDs if the SIDs are received from another segment routing implementation and the SIDs are based on the RTM active route selection.
- When SR OS does not allow assigning different SID indexes or labels to different routes of the same prefix within the same IGP instance, the router resolves only one of the duplicate SIDs if the SIDs are received from another segment routing implementation and the SIDs are based on the RTM active route selection.

The selected SID is used for ECMP resolution to all neighbors. If the route is inter-area and the conflicting SIDs are advertised by different ABRs, ECMP toward all ABRs uses the selected SID.

2.1.6.2.3 Checking for SID errors before programming the ILM and NHLFE

If any of the following conditions are true, the router logs a trap, generates a syslog error message, and does not program the ILM and NHLFE for the prefix SID:

- The received prefix SID index falls outside of the locally configured SID range.
- One or more resolved ECMP next-hops for a received prefix SID did not advertise the SR Capability sub-TLV.
- The received prefix SID index falls outside the advertised SID range of one or more resolved ECMP next-hops.

2.1.6.2.4 Programming ILM/NHLFE for duplicate prefix-SID indexes/labels for different prefixes

The router can perform the following variations of this procedure:

- For received duplicate prefix-SID indexes or labels for different prefixes within the same IGP instance, the router:
 - programs the ILM/NHLFE for the first prefix-SID index or label
 - logs a trap and generates a syslog error message
 - does not program the subsequent prefix-SID index or label in the datapath
- For received duplicate prefix-SID indexes or labels for different prefixes across IGP instances, there are two options.
 - In the global SID index range mode of operation, the resulting ILM label value is the same across the IGP instances. The router:
 - programs ILM/NHLFE for the prefix of the winning IGP instance based on the RTM route-type preference
 - logs a trap and generates a syslog error message
 - does not program the subsequent prefix SIDs in the datapath
 - In the per-instance SID index range mode of operation, the resulting ILM label has different values across the IGP instances. The router programs ILM/NHLFE for each prefix as expected.

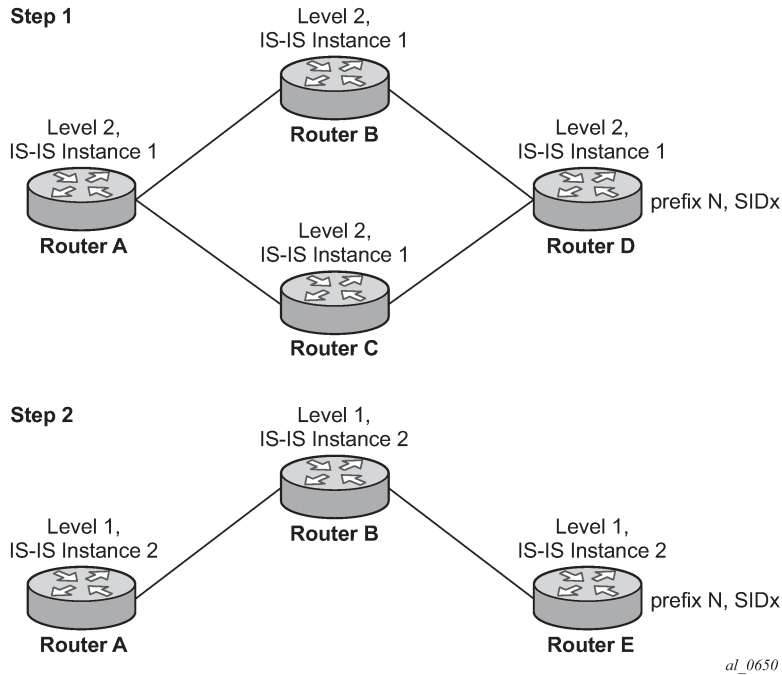
2.1.6.2.5 Programming ILM/NHLFE for the same prefix across IGP instances

In global SID index range mode of operation, the resulting ILM label value is the same across the IGP instances. The router programs ILM/NHLFE for the prefix of the winning IGP instance based on the RTM route-type preference. The router logs a trap and generates a syslog error message, and does not program the other prefix SIDs in the datapath.

In the per-instance SID index range mode of operation, the resulting ILM label has different values across the IGP instances. The router programs ILM/NHLFE for each prefix as expected.

The following figure shows an IS-IS example of handling in case of a global SID index range.

Figure 3: Handling of the same prefix and SID in different IS-IS instances



Assume the following route-type preference in RTM and tunnel-type preference in TTM are configured:

- ROUTE_PREF_ISIS_L1_INTER (RTM) 15
- ROUTE_PREF_ISIS_L2_INTER (RTM) 18
- ROUTE_PREF_ISIS_TTM 10



Note: The TTM tunnel-type preference is not used by the SR module. It is put in the TTM and is used by other applications, such as VPRN auto-bind and BGP shortcut, to select a TTM tunnel.

1. Router A performs the following resolution within the single level 2, IS-IS instance 1. All metrics are the same and ECMP = 2.
 - For prefix N, the RTM entry is:
 - prefix N
 - nhop1 = B
 - nhop2 = C
 - preference 18
 - For prefix N, the SR tunnel TTM entry is:
 - tunnel-id 1: prefix N-SIDx
 - nhop1 = B
 - nhop2 = C
 - tunl-pref 10
2. Add Level 1, IS-IS instance 2 in the same configuration, but in routers A, B, and E only.

- For prefix N, the RTM entry is:
 - prefix N
 - nhop1 = B
 - preference 15The RTM prefers L1 route over L2 route.
- For prefix N, there is one SR tunnel entry for L2 in TTM:
 - tunnel-id 1: prefix N-SIDx
 - nhop1 = B
 - nhop2 = C
 - tunl-pref 10

2.1.6.2.6 Handling ILM resource exhaustion while assigning a SID index/label

If the system exhausted an ILM resource while assigning a SID index/label to a local loopback interface, then index allocation fails and an error is displayed in the CLI. The router logs a trap and generates a syslog error message.

2.1.6.2.7 Handling ILM, NHLFE, or other IOM or CPM resource exhaustion while resolving or programming a SID index/label

If the system exhausted an ILM, NHLFE, or any other IOM or CPM resource while resolving and programming a received prefix SID or programming a local adjacency SID, the following occurs:

1. The IGP instance goes into overload and a trap and syslog error message are generated.
2. The segment routing module deletes the tunnel.

The user must manually clear the IGP overload condition after freeing resources. After the IGP is brought back up, it attempts to program all tunnels that previously failed the programming operation at the next SPF.

2.1.7 Segment routing tunnel management

The segment routing module adds a shortest path SR tunnel entry to TTM for each resolved remote node SID prefix and programs the datapath with the corresponding LTN with the push operation pointing to the primary and LFA backup NHLFEs. The LFA backup next hop for a prefix that was advertised with a node SID is only computed if the **loopfree-alternates** option is enabled in the IS-IS or OSPF instance. The resulting SR tunnel that is populated in TTM is automatically protected with FRR when an LFA backup next hop exists for the prefix of the node SID.

With ECMP, a maximum of 32 primary next-hops (NHLFEs) are programmed for the same tunnel destination for each IGP instance. ECMP and LFA next-hops are mutually exclusive, as in the current implementation.

The default preference for shortest path segment routing tunnels in the TTM is set lower than LDP tunnels but higher than BGP tunnels to allow controlled migration of customers without disrupting their current deployment when they enable segment routing.

The global default TTM preferences for the tunnel types, including the preference of both segment routing tunnels based on shortest path (referred to as SR IS-IS and SR-OSPF) is as follows:

- ROUTE_PREF_RSVP 7
- ROUTE_PREF_SR_TE 8
- ROUTE_PREF_LDP 9
- ROUTE_PREF_OSPF_TTM 10
- ROUTE_PREF_ISIS_TTM 11
- ROUTE_PREF_BGP_TTM 12
- ROUTE_PREF_GRE 255

The default value for SR IS-IS or SR-OSPF is the same, regardless of whether one or more IS-IS or OSPF instances are programming a tunnel for the same prefix. In this case, the router selects an SR tunnel based on the lowest IGP instance ID.

The TTM preference is used in the case of BGP shortcuts, VPRN auto-bind, or BGP transport tunnel when the tunnel binding commands are configured to the **any** value, which parses the TTM for tunnels in the protocol preference order. The user can use the global TTM preference or explicitly list the tunnel types to be used. When the tunnel types are listed explicitly, the TTM preference is still used to select one type over the other. In both cases, a fallback to the next preferred tunnel type is performed if the selected one fails. When a more preferred tunnel type becomes available, the system reverts to that tunnel type.

See [BGP shortcuts using segment routing tunnels](#), [BGP labeled route resolution using segment routing tunnels](#), and [Service packet forwarding with segment routing](#) for the detailed service and shortcut binding CLI.

For SR IS-IS and SR-OSPF, the user can configure the preference of each IGP instance in addition to the default values.

```
config>router>isis>segment-routing>tunnel-table-pref preference 1 to 255
config>router>ospf>segment-routing>tunnel-table-pref preference 1 to 255
```



Note: The preference of SR-TE LSP is not configurable and is the second-most preferred tunnel type after RSVP-TE. The preference of SR-TE LSP is independent of whether the SR-TE LSP was resolved in IS-IS or OSPF.

2.1.7.1 Tunnel MTU determination

The MTU of a segment routing tunnel populated into the TTM is determined in the same way as an IGP tunnel; for example, LDP LSP is based on the outgoing interface MTU minus the label stack size. Segment routing, however, supports remote LFA and TI-LFA, which can program an LFA repair tunnel by adding one or more labels.

To configure the MTU of all segment routing tunnels within each IGP instance, use the following commands:

```
config>router>isis>segment-routing>tunnel-mtu bytes bytes
config>router>ospf>segment-routing>tunnel-mtu bytes bytes
```

There is no default value for this command. If the user does not configure a segment routing tunnel MTU, the MTU, in bytes, is determined by IGP as follows:

$$\text{SR_Tunnel_MTU} = \text{MIN} \{ \text{Cfg_SR_MTU}, \text{IGP_Tunnel_MTU} - (1 + \text{frr} - \text{overhead}) \times 4 \}$$

Where:

- *Cfg_SR_MTU* is the MTU configured by the user for all segment routing tunnels within an IGP instance using the preceding CLI commands. If no value was configured by the user, the segment routing tunnel MTU is determined by the IGP interface calculation.
- *IGP_Tunnel_MTU* is the minimum of the IS-IS or OSPF interface MTU among all the ECMP paths or among the primary and LFA backup paths of this segment routing tunnel.
- *frr-overhead* is set the following parameters:
 - *value* of *ti-lfa* [**max-sr-frr-labels** labels] if **loopfree-alternates** and **ti-lfa** are enabled in this IGP instance
 - 1 if **loopfree-alternates** and **remote-lfa** are enabled but **ti-lfa** is disabled in this IGP instance
 - otherwise, it is set to 0

The SR tunnel MTU is dynamically updated anytime any of the parameters used in its calculation change. This includes when the set of the tunnel next-hops changes or the user changes the configured SR MTU or interface MTU value.



Note: The calculated SR tunnel MTU is used to determine an SDP MTU and to check the Layer 2 service MTU. When fragmenting IP packets forwarded in GRT or in a VPRN over an SR shortest path tunnel, the datapath always deducts the worst-case MTU (5 labels or 6 labels if hash label feature is enabled) from the outgoing interface MTU when deciding whether to fragment the packet. In this case, the above formula is not used.

2.1.8 Segment routing local block

Some labels that are provisioned through CLI or a management interface must be allocated from the Segment Routing Local Block (SRLB). The SRLB is a reserved label block configured under **config>router>mpls-labels**. See the *7705 SAR Gen 2 MPLS Guide* for more information about reserved label blocks.

The label block to use is specified by the **srlb** command under IS-IS or OSPF:

```
config>router>isis>segment-routing
[no] srlb reserved-label-block-name

config>router>ospf> segment-routing
[no] srlb reserved-label-block-name
```

Provisioned labels for adjacency SIDs and adjacency SID sets must be allocated from the configured SRLB. The request is rejected if any of the following are true:

- no SRLB is specified
- the requested label does not fall within the SRLB
- the label is already allocated

2.1.8.1 Bundling adjacencies in adjacency sets

An adjacency set is a bundle of adjacencies, represented by a common adjacency SID for the bundled set. It enables, for example, a path for an SR-TE LSP through a network to be specified while allowing the local node to spray packets across the set of links identified by a single adjacency SID.

SR OS supports both parallel adjacency sets (for example, those where adjacencies originating on one node terminate on a second, common node), and the ability to associate multiple interfaces on a specified node, irrespective of whether the far end of the respective links of those interfaces terminate on the same node.

An adjacency set is created under IS-IS or OSPF using the following CLI commands:

```
config
router
  isis | ospf
    segment-routing
      [no] adjacency-set id
      family [ipv4 | ipv6]
      parallel [no-advertise]
      no parallel
      exit
...
.
  exit
exit
config
router
  ospf
    segment-routing
      [no] adjacency-set id
      parallel [no-advertise]
      no parallel
      exit
...
.
  exit
```

The **adjacency-set** *id* command specifies an adjacency set, where *id* is an unsigned integer from 0 to 4294967295.

In IS-IS, each adjacency set is assigned an address family, IPv4 or IPv6. The family command for IS-IS indicates the address family of the adjacency set. For OSPF, the address family of the adjacency set is implied by the OSPF version and the instance.

The **parallel** command indicates that all members of the adjacency set must terminate on the same neighboring node. When the **parallel** command is configured, the system generates a trap message if a user attempts to add an adjacency terminating on a neighboring node that differs from the existing members of the adjacency set. See [Associating an interface with an adjacency set](#) for details about how to add interfaces to an adjacency set. The system stops advertising the adjacency set and deprograms it from TTM. The **parallel** command is enabled by default.

By default, parallel adjacency sets are advertised in the IGP. The **no-advertise** option prevents a parallel adjacency set from being advertised in the IGP; it is only advertised if the **parallel** command is configured. To prevent issues in the case of ECMP if a non-parallel adjacency set is used, an external controller may be needed to coordinate the label sets for SIDs at all downstream nodes. As a result, non-parallel adjacency sets are not advertised in the IGP. The label stack below the adjacency set label must be valid at any downstream node that exposes it, even though it is sprayed over multiple downstream next-hops.

Parallel adjacency sets are programmed in TTM (unless there is an erroneous configuration of a non-parallel adjacency). Non-parallel adjacency sets are not added to TTM or RTM, meaning they cannot be used as a hop at the originating node. Parallel adjacency sets that are advertised are included in the link-state database and TE database, but non-parallel adjacency sets are not included because they are not advertised.

An adjacency set with only one next hop is also advertised as an individual adjacency SID with the S flag set. However, the system does not calculate a backup for an adjacency set even if it has only one next hop.

2.1.8.1.1 Associating an interface with an adjacency set

IS-IS or OSPF interfaces are associated with one or more adjacency sets using the following CLI commands. Both numbered and unnumbered interfaces can be assigned to the same adjacency set.

```
config
router
  isis
    interface
      [no] adjacency-set id
      [no] adjacency-set id
      [no] adjacency-set id
config
router
  ospf
    area
      interface
        [no] adjacency-set id
        [no] adjacency-set id
        [no] adjacency-set id
```

If an interface is assigned to an adjacency set, then a common adjacency SID value is advertised for every interface in the set, in addition to the adjacency SID corresponding to the IPv4 and or IPv6 adjacency for the interface. Each IS-IS or OSPF advertisement therefore contains two adjacency SID TLVs for an address family:

- an adjacency SID for the interface (a locally-unique value)
- an adjacency SID TLV for the adjacency set

This TLV is distinguished by having the S-bit (IS-IS) or G-bit (OSPF) in the flags field set to 1. Its value is the same as other adjacency SIDs in the set at that node.

By default, both the adjacency SID for an interface and the adjacency SID for a set are dynamically allocated by the system. However, it is possible for the user to configure an alternate, static value for the SID; see [Provisioning adjacency SID values for an adjacency set](#) for more information.

A maximum of 32 interfaces can be bound to a common adjacency set. Configuring more than 32 interfaces is blocked by the system and a CLI error is generated.

Only point-to-point interfaces can be assigned to an adjacency set.

If a user attempts to assign an IES interface to an adjacency set, the system generates a CLI warning and segment routing does not program the association.

The IGP blocks the configuration of an adjacency set under an interface when the adjacency set has not yet been created under segment-routing.

In IS-IS, it is possible to add Layer 1, Layer 2, or a mix of Layer 1 and Layer 2 adjacencies to the same adjacency set.

2.1.8.1.2 Provisioning adjacency SID values for an adjacency set

For an adjacency set, static values are configured using the **sid** CLI command, as follows:

```
config>router>isis>segment-routing
  [no] adjacency-set id
    family [ipv4 | ipv6]
    [no] sid label value
    parallel [no-advertise]
    no parallel
    exit
  [no] adjacency-set id
    family [ipv4 | ipv6]
    [no] sid label value
    parallel [no-advertise]
    no parallel
    exit
  ...

config>router>ospf>segment-routing
  [no] adjacency-set id
    [no] sid label value
    parallel [no-advertise]
    no parallel
    exit
  [no] adjacency-set id
    [no] sid label value
    parallel [no-advertise]
    no parallel
    exit
  ...
```

If **no sid** is configured, a dynamic value is allocated to the adjacency set. A user may change the dynamic value to specify a static SID value. Changing an adjacency set value from dynamic to static, or static to dynamic, may result in traffic being dropped as the ILM is reprogrammed.

The *value* must correspond to a label in the reserved label block in provisioned mode referred to by the **srlb** command. A CLI error is generated if a user attempts to configure an invalid *value*. If a label is not configured, then the label *value* is dynamically allocated by the system from the dynamic labels range. If a static adjacency set label is configured, then the system does not advertise a dynamic adjacency set label.

A static label value for an adjacency set SID is persistent. Therefore, the P-bit of the flags field in the Adjacency-SID TLV, referring to the adjacency set must be set to 1.

2.1.9 Loop-free alternates

This section describes LFA implementation and configuration.

2.1.9.1 Remote LFA with segment routing

The user enables the remote LFA next-hop calculation by the IGP LFA SPF by configuring the **remote-lfa** option in the command that enables LFA calculation:

```
config>router>isis>loopfree-alternates remote-lfa
config>router>ospf>loopfree-alternates remote-lfa
```

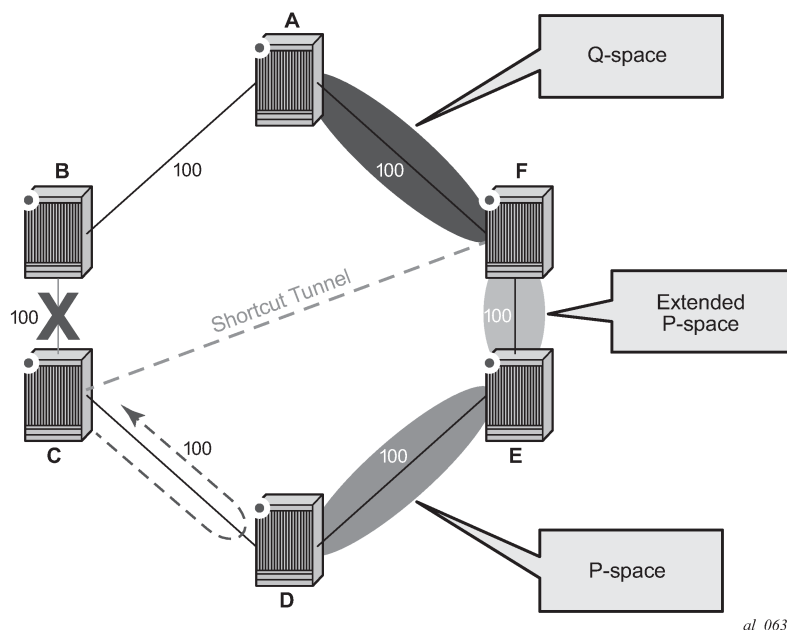
SPF performs the additional remote LFA computation following the regular LFA next-hop calculation when both of the following conditions are met.

- The **remote-lfa** option is enabled in an IGP instance.
- The LFA next-hop calculation did not result in protection for one or more prefixes resolved to a specific interface.

Remote LFA extends the protection coverage of LFA-FRR to any topology by automatically computing and establishing or tearing down shortcut tunnels or repair tunnels, to a remote LFA node, which puts the packets back into the shortest path without looping them back to the node that forwarded them over the repair tunnel. A repair tunnel can, in theory, be an RSVP LSP, an LDP-in-LDP tunnel, or an SR tunnel. In SR OS, this feature is restricted to use an SR repair tunnel to the remote LFA node.

The remote LFA algorithm for link protection is described in RFC 7490, *Remote Loop-Free Alternate (LFA) Fast Reroute (FRR)*. Unlike a typical LFA calculation, which is calculated per prefix, the LFA algorithm for link protection is a per-link LFA SPF calculation. The algorithm provides protection for all destination prefixes that share the protected link by using the neighbor on the other side of the protected link as a proxy for all of these destinations. Assume the topology in the following figure.

Figure 4: Example of remote LFA topology



When the LFA SPF in node C computes the per-prefix LFA next hop, prefixes that use link C-B as the primary next hop have no LFA next hop because of the ring topology. If node C used node link C-D as a

backup next hop, node D would loop a packet back to node C. The remote LFA then runs the following PQ algorithm, per RFC 7490.

1. Compute the extended P space of Node C with respect to link C-B. The extended P space is the set of nodes reachable from node C without any path transiting the protected link (link C-B). This computation yields nodes D, E, and F.

The determination of the extended P space by node C uses the same computation as the regular LFA by running SPF on behalf of each of the neighbors of C.



Note: RFC 7490 introduced the concept of P space, which would have excluded node F because, from the node C perspective, node C has a couple of ECMP paths, one of which goes via link C-B. However, because the remote LFA next hop is activated when link C-B fails, this rule can be relaxed and node F can be included, which then yields the extended P space.

The user can limit the search for candidate P nodes to reduce the number of SPF calculations in topologies where many eligible P nodes can exist. The following commands can be used to configure the maximum IGP cost from node C for a P node to be eligible:

- **config>router>isis>loopfree-alternates remote-lfa max-pq-cost value**
- **config>router>ospf>loopfree-alternates remote-lfa max-pq-cost value**

2. Compute the Q space of node B with respect to link C-B: the set of nodes from which the destination proxy (node B) can be reached without any path transiting the protected link (link C-B).

The Q space calculation is effectively a reverse SPF of node B. In general, one reverse SPF is run on behalf of each neighbor of C to protect all destinations resolving over the link to the neighbor. This yields nodes F and A in the example of [Figure 4: Example of remote LFA topology](#).

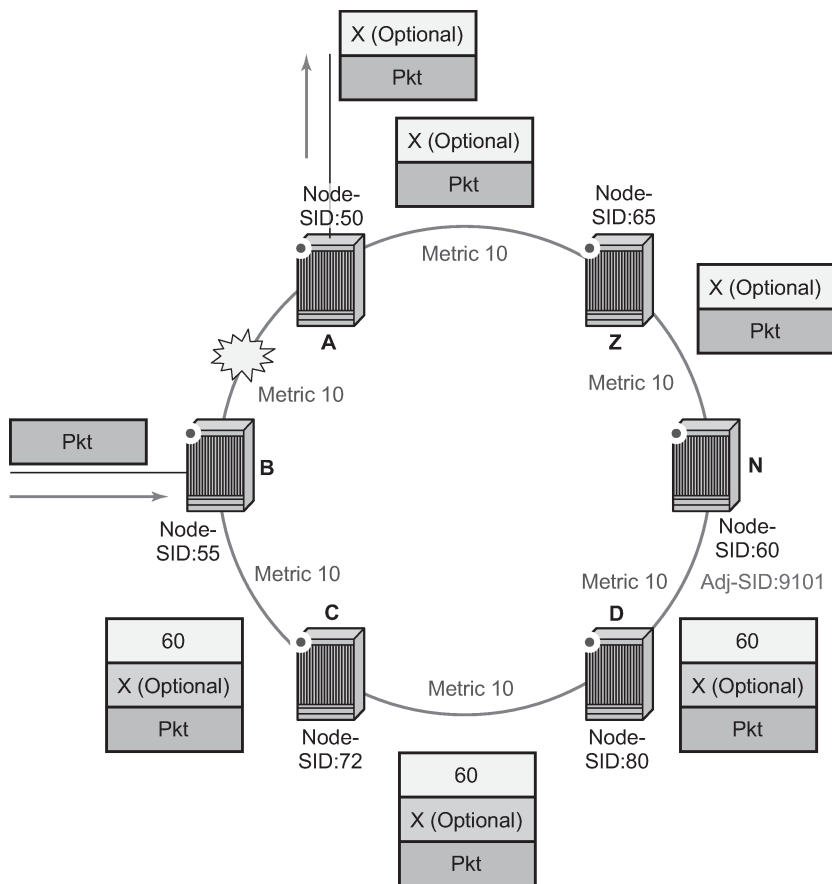
The user can limit the search for candidate Q nodes to reduce the number of SPF calculations in topologies where many eligible Q nodes can exist. The CLI commands in step 1 are also used to configure the maximum IGP cost from node C for a Q node to be eligible.

3. Select the best alternate node: this is the intersection of extended P and Q spaces. The best alternate node or PQ node is node F in the example of [Figure 4: Example of remote LFA topology](#). From node F onwards, traffic follows the IGP shortest path.

If many PQ nodes exist, the lowest IGP cost from node C is used to narrow down the selection, and if more than one PQ node remains, the node with lowest router ID is selected.

The details of the label stack encoding when the packet is forwarded over the remote LFA next hop are shown in the following figure.

Figure 5: Remote LFA next hop in segment routing



al_0648

The label corresponding to the node SID of the PQ node is pushed on top of the original label of the SID of the resolved destination prefix. If node C has resolved multiple node SIDs corresponding to different prefixes of the selected PQ node, it pushes the lowest node SID label on the packet when forwarded over the remote LFA backup next hop.

If the PQ node is also the advertising router for the resolved prefix, the label stack is compressed in the following cases, depending on the IGP:

- In IS-IS, the label stack is always reduced to a single label, which is the label of the resolved prefix owned by the PQ node.
- In OSPF, the label stack is reduced to the single label of the resolved prefix when the PQ node advertised a single node SID in this OSPF instance. If the PQ node advertised a node SID for multiple loopback interfaces within this same OSPF instance, the label stack is reduced to a single label only if the SID of the resolved prefix is the lowest SID value.

The following rules and limitations apply to the remote LFA implementation:

- If the user excludes a network IP interface from being used as an LFA next-hop using the **loopfree-alternate-exclude** command under the IS-IS or OSPF context of the interface, the interface is also excluded from being used as the outgoing interface for a remote LFA tunnel next hop.
- As with the regular LFA algorithm, the remote LFA algorithm computes a backup next hop to the ABR advertising an inter-area prefix and not to the destination prefix.

2.1.9.2 Topology-independent LFA

The Topology-Independent LFA (TI-LFA) feature improves the protection coverage of a network topology by computing and automatically instantiating a repair tunnel to a Q-node that is not in the shortest path from the computing node. The repair tunnel uses the shortest path to the P-node and a source-routed path from the P-node to the Q-node.

In addition, the TI-LFA algorithm selects the backup path that matches the post-convergence path. This helps capacity planning in the network because traffic always flows on the same path when transitioning to the FRR next hop and then on to the new primary next hop.

At a high level, the TI-LFA link protection algorithm searches for the closest Q-node to the computing node and then selects the closest P-node to this Q-node, up to a maximum number of labels. This is performed on each of the post-convergence paths to each destination node or prefix D.

When the TI-LFA feature is enabled in IS-IS, it provides a TI-LFA link-protect backup path in IS-IS MT0 for an SR IS-IS IPv4 tunnel or in MT0/MT2 for a SR IS-IS IPv6 tunnel (node SID and adjacency SID), for an IPv4/IPv6 SR-TE LSP, for an SR policy with IPv4/IPv6 endpoint, and for LDP IPv4 FEC when the LDP **fast-reroute backup-sr-tunnel** option is enabled.

2.1.9.2.1 TI-LFA configuration

Use the following command to enable TI-LFA in an IS-IS instance:

```
config>router>isis>loopfree-alternates [remote-lfa [max-pq-cost value]] [ti-lfa [max-sr-frr-labels value]]
```

The **ti-lfa** option in IS-IS provides a TI-LFA link-protect backup path in IS-IS MT0 for an SR IS-IS IPv4 tunnel and IPv6 tunnel (node SID and adjacency SID), for an IPv4/IPv6 SR-TE LSP, for an SR policy with IPv4/IPv6 endpoint, and for an LDP IPv4 FEC when the LDP **fast-reroute backup-sr-tunnel** option is enabled. For more information about the applicability of the various LFA options, see [LFA protection option applicability](#).

The **max-sr-frr-labels** parameter limits the search for the LFA backup next hop based on the value of the label as follows:

- **0**

The IGP LFA SPF restricts the search to the TI-LFA backup next hop that does not require a repair tunnel, meaning that the P-node and Q-node are the same and match a neighbor. This is also the case when both the P- and Q-nodes match the advertising router for a prefix.

- **1 to 3**

The IGP LFA SPF widens the search to include a repair tunnel to a P-node, which itself is connected to the Q-nodes with a zero to two hops for a maximum total of three labels: one node SID to the P-node and two adjacency SIDs from the P-node to the Q-node. If the P-node is a neighbor of the computing node, its node SID is compressed, meaning that up to three adjacency SIDs can separate the P- and Q-nodes.

- **2 (default)**

Corresponds to a repair tunnel to a non-adjacent P that is adjacent to the Q-node. If the P-node is a neighbor of the computing node, the node SID of the P-node is compressed, and the default value of two labels corresponds to two adjacency SIDs between the P- and Q-nodes.

If the user attempts to change the **max-sr-frr-labels** parameter to a value that results in a change to the computed FRR overhead, the IGP checks that all SR-TE LSPs can properly account for the overhead based on the configuration of the LSP **max-sr-labels** and **additional-frr-labels** parameter values; otherwise, the change is rejected.

The FRR overhead is computed by IGP and its value is set as follows:

- 0 if **segment-routing** is disabled in the IGP instance
- 0 if **segment-routing** is enabled but **remote-lfa** is disabled and **ti-lfa** is disabled
- 1 if **segment-routing** is enabled and **remote-lfa** is enabled but **ti-lfa** is disabled, or if **segment-routing** is enabled and **remote-lfa** is enabled and **ti-lfa** is enabled but **ti-lfa max-sr-frr-labels labels** is set to 0.
- the value of **ti-lfa max-sr-frr-labels labels** if **segment-routing** is enabled and **ti-lfa** is enabled, regardless if **remote-lfa** is enabled or disabled.

2.1.9.2.2 TI-LFA link-protect operation

This section describes TI-LFA protection behavior when the **loopfree-alternates** command is enabled with the **remote-lfa** and **ti-lfa** options, as described in [TI-LFA configuration](#).

2.1.9.2.2.1 LFA protection option applicability

Depending on the configured options of the **loopfree-alternates** command, the LFA SPF in an IGP instance runs the following algorithms in order.

1. The LFA SPF computes a regular LFA for each node and prefix.
In this step, a computed backup next hop satisfies any applied LFA policy. This backup next hop protects that specific prefix or node in the context of IP FRR, LDP FRR, SR FRR, SR-TE LSP and SR policy FRR.
2. The LFA SPF runs the TI-LFA algorithm if the **ti-lfa** option is enabled for all prefixes and nodes, regardless of the outcome of the first step.
If the LDP **fast-reroute backup-sr-tunnel** option is enabled, a prefix or node for which a TI-LFA backup next hop is found overrides the result from step 1 in the context of LDP FRR in SR FRR and in SR-TE FRR.
With SR FRR and SR-TE/SR policy FRR, the TI-LFA next hop protects the node-SID of that prefix and any adjacency SID terminating on the node-SID of that prefix.
The prefix or node continues to use the backup next hop found either in the context of LDP FRR (if the LDP **fast-reroute backup-sr-tunnel** option is disabled), or in the IP FRR.
3. The LFA SPF runs remote LFA only for the next hop of prefixes and nodes that remain unprotected after step 1 and step 2 if the **remote-lfa** option is enabled.
A prefix or node for which a remote LFA backup next hop is found uses it in the context of LDP FRR in SR FRR and in SR-TE FRR when the LDP **fast-reroute backup-sr-tunnel** option is enabled.

To protect an adjacency SID, the LFA selection algorithm uses the following preference order:

1. adjacency of an alternate parallel link to the same neighbor.
If more than one adjacency exists, select one as follows:
 - a. adjacency with the lowest metric

- b. adjacency to the neighbor with the lowest router ID (OSPF) or system-id (IS-IS), and the lowest metric
 - c. with the lowest interface index and the lowest router ID (OSPF) or system-id (IS-IS)
2. an ECMP next hop to a node-SID of the same neighbor that is different from the next hop of the protected adjacency.

If more than one next hop exists, select one as follows:

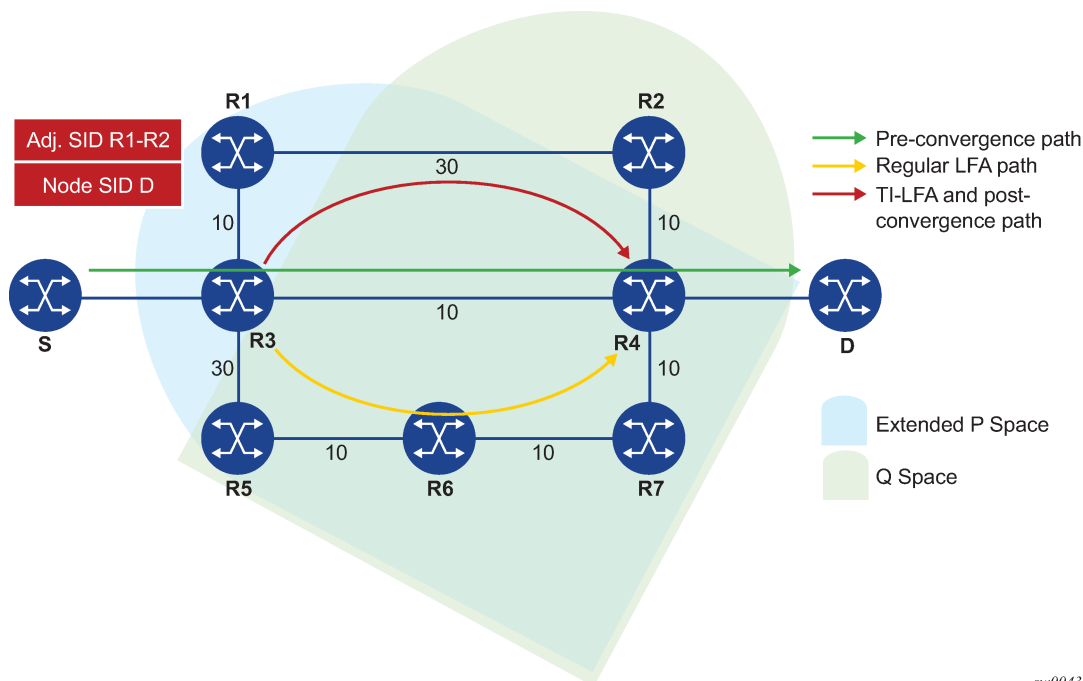
- a. next hop with the lowest metric
 - b. next hop to the neighbor with the lowest router ID (OSPF) or system-id (IS-IS) if same lowest metric
 - c. next hop to the lowest interface index if same neighbor router ID (OSPF) or system-id (IS-IS)
3. LFA backup outcome of a node SID of the same neighbor. The following is the preference order:
- a. TI-LFA backup
 - b. LFA backup
 - c. RLFA backup

2.1.9.2.2.2 TI-LFA algorithm

The TI-LFA link protection algorithm searches for the closest Q-node to the computing node and then selects the closest P-node to this Q-node, up to a number of labels corresponding to the **ti-lfa max-sr-frr-labels** *label/s* value, on each of the post-convergence paths to each destination node or prefix D.

The following figure shows a topology where router R3 computes a TI-LFA next hop for protecting link R3-R4.

Figure 6: Selecting link-protect TI-LFA backup path



sw0043

Applying the topology in the preceding figure, router R3 computes the link-protected TI-LFA backup path in the following order.

1. The router computes the post-convergence SPF on the topology without the protected link.

In the preceding figure, R3 finds a single post-convergence path to destination D via R1.



Note: The post-convergence SPF does not include IGP shortcut.

2. The router computes the extended P-Space of R3 with respect to protected link R3-R4 on the post-convergence paths.

This is the set of nodes Y_i in the post-convergence paths that are reachable from R3 neighbors without any path transiting the protected link R3-R4.

R3 computes an LFA SPF rooted at each of its neighbors within the post-convergence paths, that is, R1, using the following equation:

$$\text{Distance_opt}(R1, Y_i) < \text{Distance_opt}(R1, R3) + \text{Distance_opt}(R3, Y_i)$$

Where, $\text{Distance_opt}(A,B)$ is the shortest distance between A and B. The extended P-space calculation yields only node R1.

3. The router computes the Q-space of R3 with respect to protected link R3-R4 in the post-convergence paths.

This is the set of nodes Z_i in the post-convergence paths from which the neighbor node R4 of the protected link, acting as a proxy for all destinations D, can be reached without any path transiting the protected link R3-R4.

$$\text{Distance_opt}(Z_i, R4) < \text{Distance_opt}(Z_i, R3) + \text{Distance_opt}(R3, R4)$$

The Q-space calculation yields nodes R2 and R4.

This is the same computation of the Q-space performed by the remote LFA algorithm, except that the TI-LFA Q-space computation is performed only on the post-convergence paths.

4. For each post-convergence path, the router searches for the closest Q-node and selects the closest P-node to this Q-node, up to the number of labels corresponding to the configured **ti-lfa max-sr-fr-labels** parameter value.

The topology in [Figure 6: Selecting link-protect TI-LFA backup path](#) shows a single post-convergence path, a single P-node (R1), and that R2 is the closest of the two found Q-nodes to the P-Node.

R3 installs the repair tunnel to the P-Q set and includes the node-SID of R1 and the adjacency SID of the adjacency over link R1-R2 in the label stack. Because the P-node R1 is a neighbor of the computing node R3, the node SID of R1 is not needed and the label stack of the repair tunnel is compressed to the adjacency SID over link R1-R2 as shown in [Figure 6: Selecting link-protect TI-LFA backup path](#).

When a P-Q set is found on multiple ECMP post-convergence paths, the following selection rules are applied, in ascending order, to select a set from a single path:

- a. the lowest number of labels
- b. the next hop to the neighbor router with the lowest **router-id** (OSPF) or **system-id** (ISIS)
- c. the next hop corresponding to the Q node with the lowest **router-id** (OSPF) or **system-id** (ISIS)

If multiple links with adjacency SID exist between the selected P-node and the selected Q-node, the following rules are used for link selection:

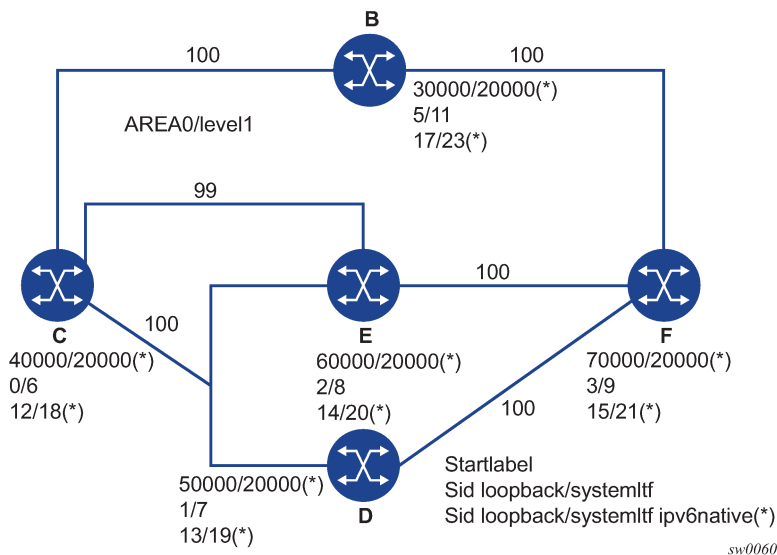
- a. the adjacency SID with the lowest metric
- b. the adjacency SID with the lowest SID value if the lowest metric is the same

2.1.9.2.2.3 TI-LFA feature interaction and limitations

The following are feature interactions and limitations of the TI-LFA link protection:

- Enabling the **ti-lfa** command in an IS-IS or OSPF instance overrides the user configuration of the **loopfree-alternate-exclude** command under the interface context in that IGP instance; that is, the TI-LFA SPF uses that interface as a backup next hop if it matches the post-convergence next hop.
- Any prefix excluded from LFA protection using the **loopfree-alternates exclude prefix-policy** command under the IGP instance context is also excluded from TI-LFA.
- Because the post-convergence SPF does not use paths transiting on a node in IS-IS overload, the TI-LFA backup path automatically does not transit on such a node.
- IES interfaces are skipped in TI-LFA computation because they do not support segment routing with MPLS encapsulation. If the only found TI-LFA backup next hop matches an IES interface, IGP treats this as if there were no TI-LFA backup paths and falls back to using either a remote LFA or regular LFA backup path in accordance with the selection rules described in [LFA protection option applicability](#).
- The TI-LFA feature provides link-protection only. Therefore, if the protected link is a broadcast interface, the TI-LFA algorithm only guarantees protection of that link and not of the pseudonode (PN) corresponding to that shared subnet. That is, if the PN is in the post-convergence path, the TI-LFA backup path may still traverse again the PN. For example, node E in [Figure 7: TI-LFA backup path via a pseudo-node](#) computes a TI-LFA backup path to destination D via E-C-PN-D because it is the post-convergence path when excluding link E-PN from the topology. This TI-LFA backup does not protect against the failure of the PN.

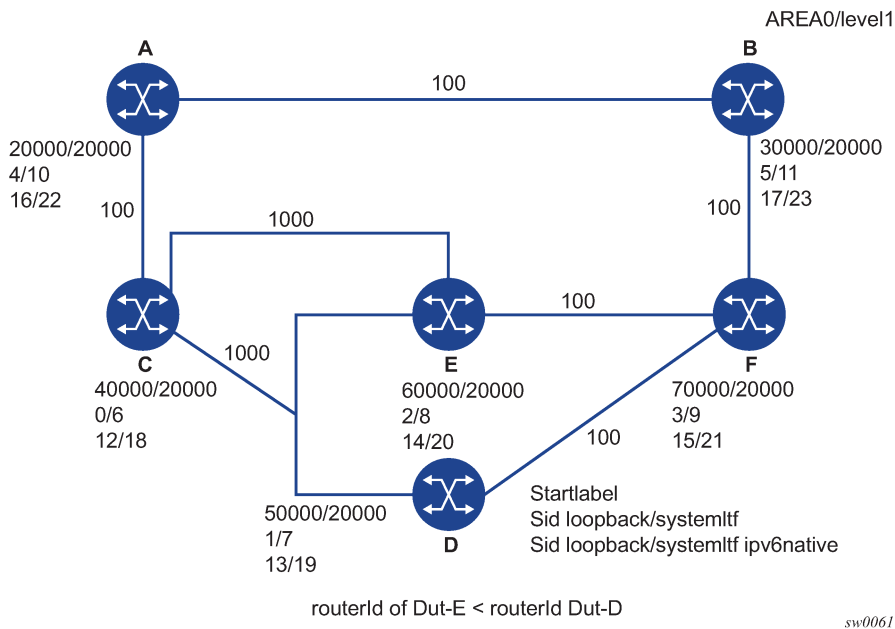
Figure 7: TI-LFA backup path via a pseudo-node



- When the computing router selects an adjacency SID among a set of parallel adjacencies between the P- and Q-nodes, the selection rules in [step 4 of TI-LFA algorithm](#). However, these rules may not yield the same interface the P-node would have selected in its post-convergence SPF, because the latter is based on the lowest value of the locally managed interface index.

For example, node A in [Figure 8: Parallel adjacencies between P and Q nodes](#) computes the link-protect TI-LFA backup path for destination node E as path A-C-E, where C is the P-node and E is the Q-node and destination. Node C has a pair of adjacency SIDs with the same metric to E. Node A selects the adjacency over the P2P link C-E because it has the lowest SID value, but node C may select the interface C-PN in its post-convergence path calculation, if that interface has a lower interface index than P2P link C-E.

Figure 8: Parallel adjacencies between P and Q nodes



- When a node SID is advertised by multiple routers (anycast SID), the TI-LFA algorithm on a router that resolves the prefix of this SID computes the backup next hop toward a single node owner of the prefix, based on the rules for prefix and SID ECMP next-hop selection.

2.1.9.2.3 Datapath support

The TI-LFA repair tunnel can have a maximum of three additional labels pushed in addition to the label of the destination node or prefix. The user can set a lower maximum value for the additional FRR labels by configuring the **ti-lfa max-sr-frr-labels labels** CLI option. The default value is 2.

The datapath models the backup path like a SR-TE LSP and therefore uses a super-NHLFE pointing to the NHLFE of the first hop in the repair tunnel. That first hop corresponds to either an adjacency SID or a node SID of the P node.

There is the special case where the P node is adjacent to the node computing the TI-LFA backup, and the Q node is the same as the P node or adjacent to the P node. In this case, the datapath at the computing router pushes either zero labels or one label for the adjacency SID between the P and Q nodes. The

backup path uses a regular NHLFE in this case, as for base LFA or remote LFA features. [Figure 6: Selecting link-protect TI-LFA backup path](#) shows an example of a single label in the backup NHLFE.

2.1.9.3 Node protection support in TI-LFA and remote LFA

This feature extends the remote LFA and TI-LFA features by adding support for node protection. The extensions are additions to the original link-protect LFA SPF algorithm.

When node protection is enabled, the router prefers a node-protect over a link-protect repair tunnel for a prefix if both are found in the remote LFA or TI-LFA SPF computations. This feature protects against the failure of a downstream node in the path of the prefix of a node SID except for the node owner of the node SID.

2.1.9.3.1 Node protection in TI-LFA and remote LFA configuration

Use the following CLI commands to configure the remote LFA and TI-LFA node protection feature.

```
configure
- router
  - [no] isis
    - [no] loopfree-alternates
      - [no] remote-lfa [max-pq-cost 0 to 4294967295, default=4261412864]
        - [no] node-protect [max-pq-nodes 1 to 32, default=16]
      - [no] ti-lfa [max-sr-frr-labels 0 to 3, default=2]
        - [no] node-protect
      - exclude
        - [no] prefix-policy prefix-policy [prefix-policy...(up to 5 max)]
      - exit
    - exit
```

The CLI commands enable the node-protect calculation to both Remote LFA (**node-protect** [**max-pq-nodes** <1 to 32, default=16>]) and TI-LFA (**node-protect**).

If the **node-protect** command is enabled, the router prefers a node-protect over a link-protect repair tunnel for a prefix if both are found in the Remote LFA or TI-LFA SPF computations. The SPF computations may only find a link-protect repair tunnel for prefixes owned by the protected node.

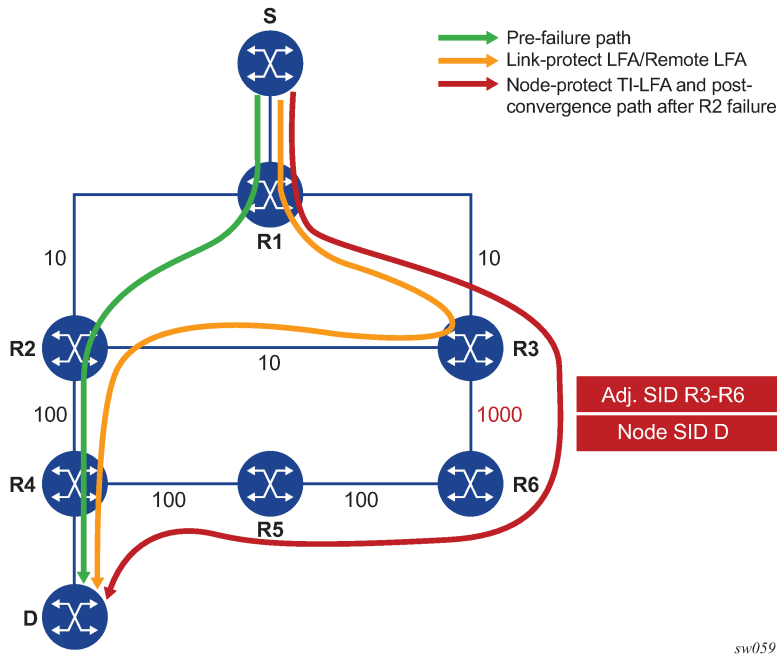
The **max-pq-nodes** parameter in the **remote-lfa** command controls the maximum number of candidate PQ-nodes found in the LFA SPFs for which the node protection check is performed. In the node protection condition, the router must run the original link-protect remote LFA algorithm plus one extra forward SPF on behalf of each PQ-node found, potentially after applying the **max-pq-cost** parameter. This checks the path from the PQ-node to the destination to ensure that the path does not traverse the protected node. Setting the **max-pq-nodes** parameter to a lower value means the LFA SPFs use less computation time and fewer resources, but may result in not finding a node-protect repair tunnel. The default value is 16. For more information, see [Remote LFA node-protect operation](#).

2.1.9.3.2 TI-LFA node-protect operation

SR OS supports the node-protect extensions to the TI-LFA algorithm as described in *draft-ietf-rtgwg-segment-routing-ti-lfa-01*.

The following figure shows a simple topology to illustrate the node-protect operation as described in [TI-LFA algorithm](#).

Figure 9: Application of the TI-LFA algorithm for node protection



The main change as a result of the node-protect extension is that the algorithm protects a node instead of a link.

Using the topology in the preceding figure, the node protection computation is performed in the following sequence:

1. The router computes the post-convergence SPF on the topology without the protected node.
In Figure 9: Application of the TI-LFA algorithm for node protection, R1 computes TI-LFA on the topology without the protected node R2 and finds a single post-convergence path to destination D via R3 and R6.

Prefixes owned by all other nodes in the topology have a post-convergence path via R3 and R6 except for prefixes owned by node R2. The latter uses the link R3-R2 and they can only benefit from link protection.

2. The router computes the extended P-Space of R1 with respect to protected node R2 on the post-convergence paths.

This is the set of nodes Y_i in the post-convergence paths that are reachable from R1 neighbors, other than protected node R2, without any path transiting the protected node R2.

R1 computes an LFA SPF rooted at each of its neighbors within the post-convergence paths. For example, R1 uses the following calculation to compute the LFA SPF for R3:

$$\text{Distance_opt}(R3, Y_i) < \text{Distance_opt}(R3, R2) + \text{Distance_opt}(R2, Y_i)$$

Where:

Distance_opt(A,B) is the shortest distance between A and B.

The extended P-space calculation yields node R3 only.

3. The router computes the Q-space of R1 with respect to protected link R1-R2 on the post-convergence paths.

This is the set of nodes Z_i in the post-convergence paths from which node R2 can be reached without any path transiting the protected link R1-R2, using the following equation:

$$\text{Distance_opt}(Z_i, R2) < \text{Distance_opt}(Z_i, R1) + \text{Distance_opt}(R1, R2)$$

The reverse SPF for the Q-space calculation is the same as in the link-protect algorithm and uses the protected node R2 as the proxy for all destination prefixes. To compute the Q space with respect to the protected node R2 instead of link R1-R2, a reverse SPF would have to be performed for each destination D which is very costly and not scalable. However, this means the path from the Q-node to the destination D or the path from the P-node to the Q-node is not guaranteed to avoid the protected node R2. The intersection of the Q-space with post-convergence path is modified in the next step to mitigate this risk.

This step yields nodes R3, R4, R5, and R6.

4. For each post-convergence path, the router searches for the closest Q-node to destination D and selects the closest P-node to this Q-node, up to the number of labels corresponding to the configured **ti-lfa max-sr-frr-labels** parameter value.

This step yields the following P-Q sets, depending on the value of the **max-sr-frr-labels** parameter:

- **max-sr-frr-labels=0**

R3 is the closest Q-node to the destination D and R3 is the only P-node. This case results in link protection via PQ-node R3.

- **max-sr-frr-labels=1**

R6 is the closest Q-node to the destination D and R3 is the only P-node. The repair tunnel for this case uses the SID of the adjacency over link R3-R6 and is shown in [Figure 9: Application of the TI-LFA algorithm for node protection](#).

- **max-sr-frr-labels=2**

R5 is the closest Q-node to the destination D and R3 is the only P-node. The repair tunnel for this case uses the SIDs of the adjacencies over links R3-R6 and R6-R5.

- **max-sr-frr-labels=3**

R4 is the closest Q-node to the destination D and R3 is the only P-node. The repair tunnel for this case uses the SIDs of the adjacencies over links R3-R6, R6-R5, and R5-R4.

This step of the algorithm is modified from link protection, which prefers Q-nodes that are the closest to the computing router R1. This is to minimize the probability that the path from the Q-node to the destination D, or the path from the P-node to the Q-node, goes via the protected node R2 as described in step 2. However, there is still a probability that the found P-Q set achieves link protection only.

5. Select the P-Q Set.

If a candidate P-Q set is found on each of the multiple ECMP post-convergence paths in step 4, the following selection rules are applied in ascending order to select a single set:

- a. the lowest number of labels
- b. the lowest next-hop router ID
- c. the lowest interface index if the same as the next-hop router ID

If multiple parallel links with adjacency SID exist between the P- and Q-nodes of the selected P-Q set, the following rules are used to select one of them:

- a. the adjacency SID with lowest metric
- b. the adjacency SID with the lowest SID value, if the same as the lowest metric

For each destination prefix D, R1 programs the TI-LFA repair tunnel (**max-sr-frr-labels=1**):

- For prefixes other than those owned by node R2 and R3, R1 programs a node-protect repair tunnel to the P-Q pair R3-R6 by pushing the SID of adjacency R3-R6 on top of the SID for destination D and programming a next hop of R3.
- For prefixes owned by node R2, R1 runs the link-protect TI-LFA algorithm and programs a simple link-protect repair tunnel, which consists of a backup next hop of R3 and pushes no additional label on top of the SID for the destination prefix.
- Prefixes owned by node R3 are not impacted by the failure of R2 because their primary next hop is R3.

2.1.9.3.3 Remote LFA node-protect operation

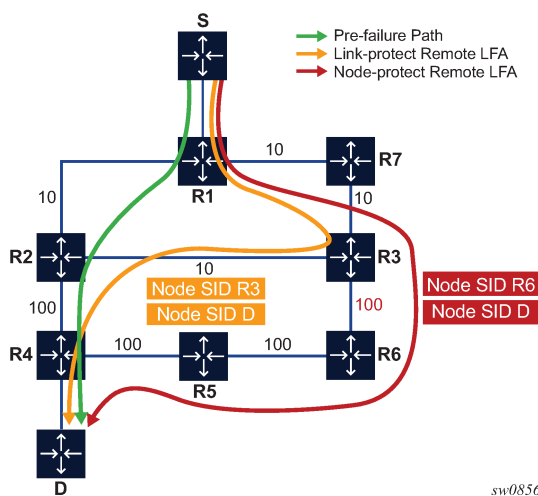
SR OS supports the node-protect extensions to the Remote LFA algorithm, as described in RFC 8102.

Remote LFA follows a similar algorithm as TI-LFA but does not limit the scope of the calculation of the extended P-Space and of the Q-Space to the post-convergence paths.

Remote LFA adds an extra forward SPF on behalf of the PQ node to ensure that, for each destination, the selected PQ-node does not use a path via the protected node.

The following figure shows a slightly modified topology from that in [TI-LFA feature interaction and limitations](#). A new node R7 is added to the top ring and the metric for link R3-R6 is modified to 100.

Figure 10: Application of the remote LFA algorithm for node protection



Using the topology in the preceding figure, the node-protect remote LFA algorithm computation is performed in the following sequence:

1. Compute extended P-Space of R1 with respect to protected node R2.

This is the set of nodes Y_i which are reachable from R1 neighbors, other than protected node R2, without any path transiting the protected node R2.

R1 computes a LFA SPF rooted at each of its neighbors, in this case, R7, using the following equation:

$$\text{Distance_opt}(R7, Y_i) < \text{Distance_opt}(R7, R2) + \text{Distance_opt}(R2, Y_i)$$

Where $\text{Distance_opt}(A,B)$ is the shortest distance between A and B.

Nodes R7, R3 and R6 satisfy this inequality.

2. Compute Q-space of R1 with respect to protected link R1-R2.

This is the set of nodes Z_i from which node R2 can be reached without any path transiting the protected link R1-R2, using the following equation.

$$\text{Distance_opt}(Z_i, R2) < \text{Distance_opt}(Z_i, R1) + \text{Distance_opt}(R1, R2)$$

The reverse SPF for the Q-space calculation is the same as in the remote LFA link-protect algorithm and uses the protected node R2 as the proxy for all destination prefixes.

This step yields nodes R3, R4, R5, and R6.

Therefore, the candidate PQ nodes after this step are nodes R3 and R6.

3. For each PQ node found, run a forward SPF to each destination D.

This step is required to select only the subset of PQ-nodes that do not traverse protected node R2.

$$\text{Distance_opt}(PQ_i, D) < \text{Distance_opt}(PQ_i, R2) + \text{Distance_opt}(R2, D)$$

Of the candidates PQ nodes R3 and R6, only PQ node R6 satisfies this inequality.

This step of the algorithm is applied to the subset of candidate PQ-nodes out of steps 1 and 2 and to which the **max-pq-cost** parameter was already applied. This subset is further reduced in this step by retaining the candidate PQ-nodes that provide the highest coverage among all protected nodes in the topology, and the number of which does not exceed the value of the **max-pq-nodes** parameter.

In case of multiple candidate PQ nodes out of this step, the detailed selection rules of a single PQ-node from the candidate list is provided in step 4.

4. Select a PQ-Node.

If multiple PQ nodes satisfy the criteria in all the above steps, then R1 further selects the PQ node as follows.

- a. R1 selects the lowest IGP cost from R1.
- b. If more than one PQ-nodes remains, R1 selects the PQ-node reachable via the neighbor with the lowest router ID (OSPF) or system ID (IS-IS).
- c. If more than one PQ-node remains, R1 selects the PQ node with the lowest router ID (OSPF) or system ID (IS-IS).

For each destination prefix D, R1 programs the remote LFA backup path as follows:

- For prefixes of R5, R4 or downstream of R4, R1 programs a node-protect remote LFA repair tunnel to the PQ node R6 by pushing the SID of node R6 on top of the SID for destination D and programming a next hop of R7.
- For prefixes owned by node R2, R1 runs the link-protect remote LFA algorithm and programs a simple link-protect repair tunnel which consists of a backup next hop of R7 and pushing the SID of PQ node R3 on top of the SID for the destination prefix D.

- Prefixes owned by nodes R7, R3, and R6 are not impacted by the failure of R2 because their primary next hop is R7.

2.1.9.3.4 TI-LFA and remote LFA node protection feature interaction and limitations

The order of activation of the various LFA types on a per prefix basis is as follows: TI-LFA, followed by base LFA, followed by remote LFA. See [LFA protection option applicability](#) for more information about the order of activation.

Node protection is enabled for TI-LFA and remote LFA separately. The base LFA prefers node protection over link protection.

The order of activation of the LFA types supersedes the protection type (node versus link). Consequently, a prefix can be programmed with a link-protect backup next hop by the more preferred LFA type. For example, a prefix is programmed with the only link-protect backup next hop found by the base LFA when a node-protect remote LFA next hop exists.

2.1.9.4 LFA policies

This section describes the application of LFA policies.

2.1.9.4.1 Application of LFA policy to a segment routing node SID tunnel

When a route next-hop policy template is applied to an interface, the LFA backup selection algorithm is extended to also apply to IPv4/IPv6 SR IS-IS, and IPv4 SR-OSPF node-SID tunnels in which a primary next hop is reachable using that interface. The extension applies to the following LFA methods: base LFA, remote LFA (RLFA), and Topology-Independent LFA (TI-LFA).

The following general rules apply across all LFA methods:

- The LFA policy constraints admin-group (**include-group** and **exclude-group**) and SRLG (**srlg-enable**) are only checked against the outgoing interface used by the LFA, RLFA, or TI-LFA backup path.
- The LFA policy parameter **protection-type {link | node}**, which controls the preference among link and node protection backup types, applies to all LFA methods.

The base LFA automatically computes both protection backup path types but, on a prefix basis, by default prefers to enforce the node-protect over the link-protect backup next hop.

By default, RLFA and TI-LFA compute only the link-protect backup path, unless the optional command **node-protect** is enabled, in which case, the preference is reversed.

For all three LFA methods, when the LFA policy enables a preference for link-protect or node-protect, the backup path is selected from the computed paths based on the configuration for the individual LFA method protection preference and the outcome (node-protect or link-protect) of the actual computation within each method. However, on a per-destination prefix basis, the post-convergence constraint of TI-LFA is selected over the LFA protection type in all cases. The selection rule uses the TI-LFA backup (if one exists), even if it is of a less-preferred protection type than the backup path computed by base LFA and RLFA.

For example, assume that an LFA policy with **protection-type=node** is applied to an IS-IS interface and the **node-protect** command is enabled in both RLFA and TI-LFA contexts in this IS-IS instance.

If TI-LFA found a link-protect backup path for the destination prefix of a SR IS-IS tunnel, it is always selected over the base LFA node-protect and RLFA node-protect backup paths.

The outcomes of LFA policy selections for specified destination prefixes of SR tunnels are described in the following tables.

Table 2: Outcome of LFA policy with protection-type=node

RLFA outcome	LFA policy protection-type=node								
	Base LFA outcome								
	none			link-protect			node-protect		
	TI-LFA outcome			TI-LFA outcome			TI-LFA outcome		
	none	link-protect	node-protect	none	link-protect	node-protect	none	link-protect	node-protect
none	none	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA
link-protect	RLFA	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA
node-protect	RLFA	TI-LFA	TI-LFA	RLFA	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA

Table 3: Outcome of LFA policy with protection-type=link

RLFA Outcome	LFA policy protection-type=link								
	Base LFA outcome								
	none			link-protect			node-protect		
	TI-LFA outcome			TI-LFA outcome			TI-LFA outcome		
	none	link-protect	node-protect	none	link-protect	node-protect	none	link-protect	node-protect
none	none	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA
link-protect	RLFA	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA
node-protect	RLFA	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA	LFA	TI-LFA	TI-LFA

- LFA policy parameter **nh-type {ip | tunnel}**, which controls preference among the backup of type IP and type tunnel (IGP shortcut), is not applicable to RLFA and TI-LFA backup paths.

However, the parameter applies if the LFA policy results in selecting a base LFA backup and the user-enabled resolution of SR-ISIS or SR-OSPF tunnel over IGP shortcut using RSVP-TE LSP.

- When configured on an interface, the route next-hop policy template applies to destination prefixes of the following types:
 - IPv4 and IPv6 SR IS-IS node SID tunnels
 - IPv4 SR-OSPF node SID tunnels
 - IPv4 tunnels where the primary next hop is reachable using that interface

The route next-hop policy template also indirectly applies to the following:

- IPv4 or IPv6 SR-TE LSPs
- IPv4 or IPv6 SR policies that use any of the previously mentioned SR tunnels as the top SID in their SID list

Finally, the LFA policy indirectly applies to IPv4 LDP FECs when the LDP **fast-reroute backup-sr-tunnel** command is enabled and the FEC is protected using an SR tunnel.

- An LFA policy applied to an interface cannot be selectively enabled or disabled per LFA method.
- As a result of these rules, no more than one backup path remains in each LFA method. In this case, the selection preference order is as follows:
 1. TI-LFA backup IP next hop or repair tunnel
 2. base LFA backup next hop

This can be of type IP (default or if **nh-type** type preference set to **ip**) or of type tunnel (**nh-type** type preference is set to **tunnel** and family SRv4 or SRv6 resolves to IGP shortcut using RSVP-TE LSP).
 3. remote LFA repair tunnel

2.1.9.4.1.1 Modifying the remote LFA selection algorithm

This section provides detailed steps for modifying the RLFA selection algorithm. The admin-group and SRLG constraints are applied to the neighbors [Ni] prior to the computation of the candidate PQ nodes.

The candidate PQ computations are run for both link protection (link S-E removed) and node protection (node E removed) because there is a need to fall back to the less preferred protection option in accordance with the value of the **protection-type** parameter in the LFA policy applied to a prefix.

Perform the following steps to modify the RLFA selection algorithm.

1. Apply the LFA policy, which is the policy that corresponds to the protected link S-E. If the **node-protect** command is enabled in RLFA, the applied LFA policy is the one corresponding to the primary next hop to the protected node.
2. Apply the following admin-group constraints to each neighbor [Ni].
 - a. Prune links that do not include one or more of the admin-groups in the **include-group** statements in the route next-hop policy template.
 - b. Prune links that belong to admin-groups that have been explicitly excluded using the **exclude-group** statement in the route next-hop policy template.
 - c. Exclude a neighbor [Ni] when it is only reachable by interfaces that violate the previous admin-group constraints.
3. Apply the following SRLG constraints to each neighbor [Ni].
 - a. Prune links that belong to the SRLGs used by the primary next hop of a destination prefix.
 - b. Exclude a neighbor [Ni] when only reachable by interfaces that violate the previous SRLG constraint.
4. Apply one of the following **protection-type** preferences:
 - Perform one of the following if **protection-type=link**.
 - If **node-protect** is disabled in RLFA, select the results of the link-protect calculation. In some cases, the computed backup might be node-protecting but shows as link-protect in the output of the **tools dump router ospf/isis sr-database** command.

- If **node-protect** is enabled in RLFA, select the results of the link-protect calculation in preference over the results of the node-protect calculation.
 - Perform one of the following if **protection-type=node**.
 - If **node-protect** is disabled in RLFA, select the results of the link-protect calculation.
 - If **node-protect** is enabled in RLFA, select the results of the node-protect calculation in preference over the results of the link-protect calculation.
5. Apply the **next-hop** type preference (not applicable to RLFA).
 6. Select the best [Ni] next hop among the remaining ones in the paths as candidate PQ nodes of prefix E (acting as proxy for destination prefix D), according to the following rules (in ascending order):
 - a. prefer the next hop, avoiding the pseudo-node (PN) used by the primary next hop
 - b. within the remaining subset, prefer the node-protect type or link-protect type according to the value of the **protection-type** option in the route next-hop policy template
 - c. within the remaining subset, select the best admin group or groups according to the preference specified in the value of the **include-group** option in the route next-hop policy template
 - d. select the [Ni] next hops corresponding to the PQ nodes with the same lowest cost (for example, the closest to RLFA computing node [S])
 - e. if more than one [Ni] next hop remains, select the next hops of the [Ni] with the lowest router ID (OSPF) or system ID (IS-IS)
 - f. if more than one [Ni] next hop remains, select the next hops of the [Ni] corresponding to the PQ node with the lowest router ID (OSPF) or system ID (IS-IS)

2.1.9.4.1.2 Modifying the TI-LFA selection algorithm

This section provides detailed steps to modify the TI-LFA selection algorithm.

The SRLG constraint is applied to the interfaces to each neighbor [Ni] prior to performing the post-convergence SPF on the topology with link S-E removed (**link-protect**) or node E removed (**node-protect**).

The admin-group constraint is applied to the outgoing interfaces of the next hops of neighbors [Ni] resulting from the post-convergence SPF computation of the destination prefix D with link S-E removed (link-protect) or node E removed (node-protect). Consequently, the number of next hops and outgoing interfaces selected by the post-convergence SPF, which is influenced by the router **ecmp** value, may violate the LFA policy constraints.

Therefore, the destination prefix may remain unprotected, or may be protected with a less-preferred next hop by TI-LFA even if another LFA policy complies or a more preferred outgoing link exists but was not selected by the post-convergence SPF. This is because the post-convergence SPF part of TI-LFA must select the same outgoing interface and next hop as the post-convergence main SPF computation performed by the node for the destination prefix

The post-convergence SPF and PQ set computations are run for both link protection (with link S-E removed) and node protection (with node E removed), because there is a need to fall back to the less-preferred protection option according to the value of parameter **protection-type** in the LFA policy applied to a prefix.

Perform the following steps to modify the TI-LFA selection algorithm.

1. Apply the LFA policy that corresponds to the protected link S-E. If the **node-protect** command is enabled in TI-LFA, the applied LFA policy corresponds to the link that is the primary next hop to the protected node.
2. Run a post-convergence SPF on the modified topology resulting from:
 - pruning the link of any outgoing interface to a neighbor N[i] that shares one or more SRLGs with the outgoing interface of the primary path of the destination prefix D
 - pruning the node of any neighbor N[i] when all outgoing interfaces to this neighbor share one or more SRLGs with the outgoing interface of the primary path of the destination prefix D
3. Apply the following admin-group constraints to each outgoing interface to a neighbor [Ni] in a post-convergence path to a destination prefix D.
 - a. Prune links that do not include one or more of the admin-groups in the **include-group** statements in the route next-hop policy template.
 - b. Prune links that belong to admin-groups that have been explicitly excluded using the **exclude-group** statement in the route next-hop policy template.
 - c. Exclude a neighbor [Ni] when it is only reachable by outgoing interfaces that violate the previous admin-group constraints.
4. Apply one of the following **protection-type** preferences:
 - Perform one of the following if **protection-type=link**.
 - If **node-protect** is disabled in TI-LFA, select the results of link-protect calculation. In some cases, the computed backup might be node-protecting but show as link-protect in the output of the **tools dump router ospf/isis sr-database** command.
 - If **node-protect** is enabled in TI-LFA, select the results of link-protect calculation in preference over the results of the node-protect calculation.
 - Perform one of the following if **protection-type=node**.
 - If **node-protect** is disabled in TI-LFA, select the results of the link-protect calculation.
 - If **node-protect** is enabled in TI-LFA, select the results of the node-protect calculation in preference over the results of the link-protect calculation.
5. Apply the **next-hop** type preference (not applicable to TI-LFA).
6. Select the best [Ni] next hop among the remaining ones in the paths as candidate PQ sets of destination D, according to the following rules (in ascending order):
 - a. prefer the next hop, avoiding the pseudo-node (PN) used by the primary next hop
 - b. within the remaining subset, prefer the node-protect type or the link-protect type according to the value of the **protection-type** option in the route next-hop policy template
 - c. within the remaining subset, select the best admin group or groups according to the preference specified in the value of the **include-group** option in the route next-hop policy template
 - d. select the [Ni] next hops corresponding to the PQ sets with the lowest number of labels
 - e. if more than one remains, select the next hops to the [Ni] with the lowest router ID (OSPF) or system ID (IS-IS)
 - f. if more than one remains, select the next hops to the [Ni] corresponding to the Q node with the lowest router ID (OSPF) or system ID (IS-IS)

2.1.9.4.2 Application of LFA policy to adjacency SID tunnel

The modifications to TI-LFA and RLFA as described in [Application of LFA policy to a segment routing node SID tunnel](#) are similarly applied to an adjacency SID tunnel.

The LFA selection algorithm for an adjacency to a neighbor uses the following preference order:

1. Adjacency of an alternate parallel link to the same neighbor, determined as follows:
 - a. apply admin-group and SRLG constraints of the LFA policy of the link of the protected adjacency
 - b. select the adjacency with best admin-groups according to the preference specified in the value of the **include-group** option in the route next-hop policy template
 - c. select the adjacency with lowest metric
 - d. select the adjacency to the neighbor with the lowest router ID (OSPF) or system ID (IS-IS), and the lowest metric
 - e. select the adjacency over the lowest interface index, and the lowest neighbor router ID (OSPF) or system ID (IS-IS)
2. ECMP next hop to a node-SID of the same neighbor, determined as follows:
 - a. apply admin-group and SRLG constraints of the LFA policy of the link of the protected adjacency
 - b. select the next hop with the best admin-groups according to the preference specified in the value of the **include-group** option in the route next-hop policy template
 - c. select the next hop with lowest metric
 - d. select the next hop to the neighbor with the lowest router ID (OSPF) or system ID (ISIS), and the lowest metric
 - e. select the next hop over the lowest interface index, and the lowest neighbor router ID (OSPF) or system ID (IS-IS)
3. LFA backup outcome of a node SID of the same neighbor:

select a LFA backup with an outgoing link that does not conflict with the LFA policy of the link of the protected adjacency



Note: If a different LFA policy was already applied in the computation of the LFA backup of the node SID of the neighbor, it is possible that some links to that node SID may have been eliminated before applying the LFA policy of the link of the protected adjacency.

2.1.9.4.3 Application of LFA policy to backup node SID tunnel

The backup node SID feature allows OSPF to use the path to an alternate ABR as an RLFA backup for forwarding packets of prefixes outside the local area or domain when the path to the primary ABR fails.

This feature reduces the label stack size by omitting the PQ node label if a regular RLFA algorithm is run.

The backup node SID algorithm consists of the following steps:

1. Perform an SPF on the modified topology with the primary ABR removed.

This action resolves the backup node SID using the path to the alternate ABR.
2. Install the ILM to use the backup node SID for transit traffic with the maximum ECMP next hops found in step 1.

3. Use the backup node SID as an RLFA backup for prefixes outside the local area or domain. This step is modified as follows to select the backup node SID by applying the LFA policy corresponding the primary next hop of these prefixes, as follows.
 - a. For each neighbor (Ni) found in step 1, use the LFA policy to select the best next-hop interface.
 - b. Among the remaining interfaces, use the LFA policy to select best (Ni) and select its interface.

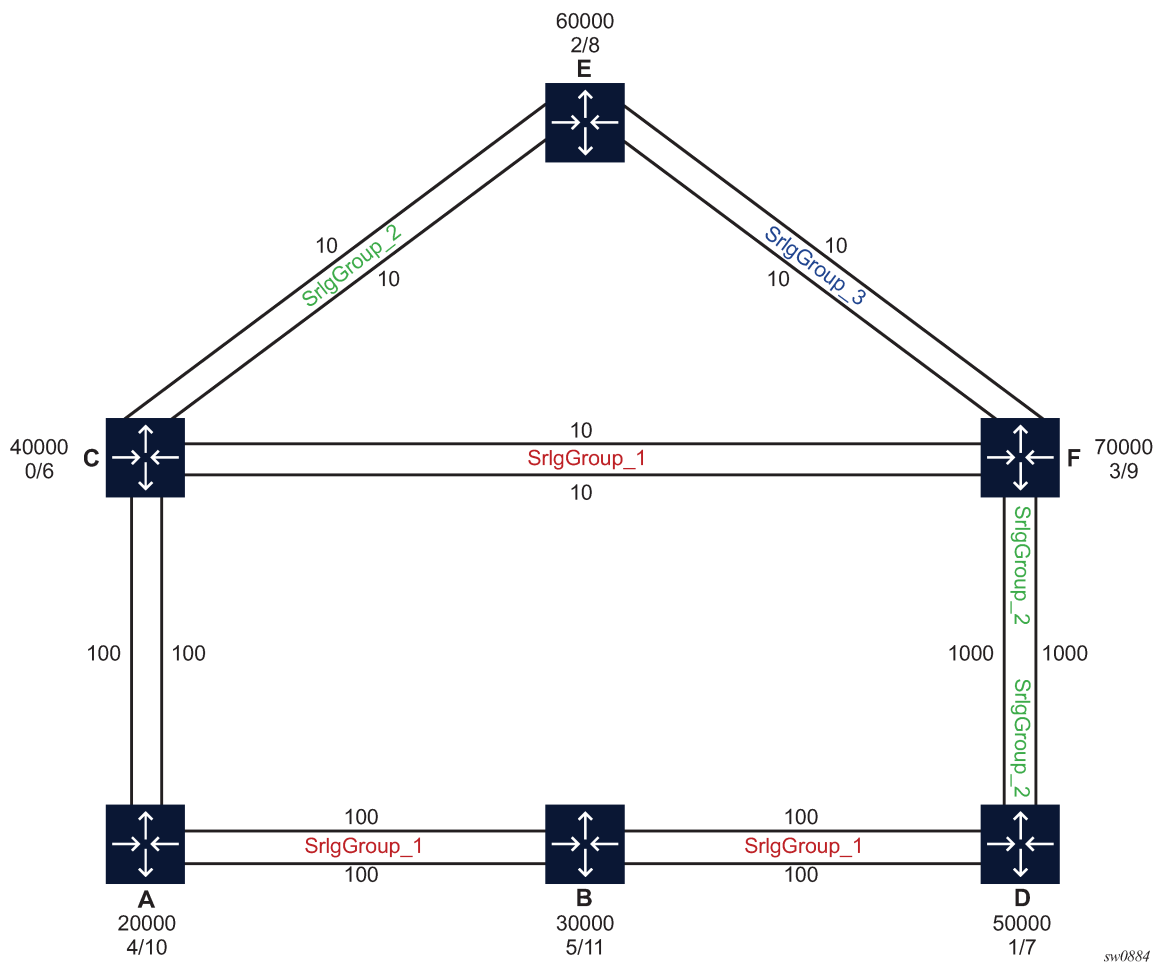


Note: A backup node SID is always preferred to a regular RLFA backup. This does not change after applying the LFA policy because the main objective of the backup node SID feature is to reduce the label stack size of the backup tunnel.

2.1.9.4.4 Configuration example of LFA policy use in remote LFA and TI-LFA

The following figure shows an example network topology that uses the OSPF routing protocol and in which the user assigns an SRLG ID to each group of OSPF links to represent fate-sharing among the links in the group. Assume the router **ecmp** value is set to 1.

Figure 11: Application of LFA policy to RLFA and TI-LFA



The user wants to enforce that the LFA backup computed and programmed by each node for a specific destination prefix avoids the SRLG ID of the primary next hop of that prefix. To that effect, the user applies an LFA policy to each link that is used as a primary next hop to reach destination prefixes.

For example, node F uses the top interface to node C as the primary next hop for the SR-OSPF tunnel to the SID of node C. The LFA policy states that the LFA backup must exclude outgoing interfaces that are members of the SRLG ID of the interface of the primary next hop. Therefore, node F must select an LFA backup that avoids SRLG ID=SrlgGroup_1.

Node F enables base LFA, remote LFA with **node-protect**, and TI-LFA with **node-protect** on the OSPF routing instance. The LFA SPF yields the following candidate LFA backup paths for the tunnel to the SID of node C.

1. Base LFA returns two backup paths: next hop over the second interface to C (cost 10) and next hop over the interface to node E (cost 20).

After applying the LFA policy, only next hop over the interface to node E (cost 20) remains. The second interface to node C is also a member of SRLG ID=SrlgGroup_1 and, therefore, the LFA next hop using it is excluded.

2. After pruning the second interface to node C, which is also a member of SRLG ID=SrlgGroup_1, TI-LFA returns a single backup path via PQ node E over the interface to node E (cost 20).

3. Remote LFA returns two backup paths, one backup path by PQ node C over the second interface to C (cost 10) and one by PQ node E over the interface to node E (cost 20).

After applying the LFA policy, only the backup path by PQ node E over the interface to node E (cost 20) remains.

4. The LFA backup paths found by all three LFA methods are only link-protecting because node C is a neighbor of node F.

5. The final outcome is the selection among the LFA methods and TI-LFA is preferred to base LFA, which is preferred to RLFA; therefore, the TI-LFA next hop over the interface to node E (cost 20) is selected and programmed by node F as the backup path for the SR-OSPF tunnel to the SID of node C.

6. The adjacency from node F to node C over first interface to node C also inherits the same LFA backup path as the node SID of C because the same LFA policy applies.

Example

The following are excerpts of the CLI configuration of node F in this specific example. The commands relevant to the LFA policy applied to link F-C are identified by arrows.

In addition, the output of show commands in node F highlights both the primary and the link-protect TI-LFA backup for both the node SID tunnel to C and the adjacency SID tunnel over the first interface to node C.

Because C is the termination for both its node SID and the adjacency SID tunnels from node F, only link protection can be provided as shown by the output of **tools>dump>router>ospf sr-database** command (field L(R)). However, the output of the same show command for the tunnel to the SID of node D indicates the TI-LFA backup over the direct interface to node D is node-protecting (field Tn(R)).

```
*A:Dut-F>config>router# info
-----
#-----
echo "IP Configuration"
#-----
if-attribute <-----
```

```

        srlg-group "SrlgGroup_1" value 1 <-----
        srlg-group "SrlgGroup_2" value 2
        srlg-group "SrlgGroup_3" value 3
    exit
    route-next-hop-policy <-----
    begin <-----
    template "templateSrlgGroup_1" <-----
        srlg-enable
    exit
    template "templateSrlgGroup_2"
        srlg-enable
    exit
    template "templateSrlgGroup_3"
        srlg-enable
    exit
    commit
exit
interface "DUTF_TO_DUTC.1.0" <-----
    address 1.0.36.6/24
    secondary 51.0.36.6/24
    port 1/1/4:1
    mac 00:00:00:00:00:06
    ipv6
        address 3ffe::100:2406/120 primary-preference 1
        address 3ffe::3300:2406/120 primary-preference 2
    exit
    if-attribute <-----
        srlg-group "SrlgGroup_1" <-----
    exit
    no shutdown
exit
interface "DUTF_TO_DUTC.2.0" <-----
    address 2.0.36.6/24
    secondary 52.0.36.6/24
    port 1/1/4:2
    mac 00:00:00:00:00:06
    ipv6
        address 3ffe::200:2406/120 primary-preference 1
        address 3ffe::3400:2406/120 primary-preference 2
    exit
    if-attribute <-----
        srlg-group "SrlgGroup_1" <-----
    exit
    no shutdown
exit
interface "DUTF_TO_DUTD.1.0"
    address 1.0.46.6/24
    secondary 51.0.46.6/24
    port 1/1/1:1
    mac 00:00:00:00:00:06
    ipv6
        address 3ffe::100:2e06/120 primary-preference 1
        address 3ffe::3300:2e06/120 primary-preference 2
    exit
    if-attribute
        srlg-group "SrlgGroup_2"
    exit
    no shutdown
exit
interface "DUTF_TO_DUTD.2.0"
    address 2.0.46.6/24
    secondary 52.0.46.6/24
    port 1/1/1:2
    mac 00:00:00:00:00:06

```

```

        ipv6
            address 3ffe::200:2e06/120 primary-preference 1
            address 3ffe::3400:2e06/120 primary-preference 2
        exit
        if-attribute
            srlg-group "SrlgGroup_2"
        exit
        no shutdown
    exit
interface "DUTF_TO_DUTE.1.0"          <-----
    address 1.0.56.6/24
    secondary 51.0.56.6/24
    port 1/1/2:1
    mac 00:00:00:00:00:06
    ipv6
        address 3ffe::100:3806/120 primary-preference 1
        address 3ffe::3300:3806/120 primary-preference 2
    exit
    if-attribute                        <-----
        srlg-group "SrlgGroup_3"      <-----
    exit
    no shutdown
exit
interface "DUTF_TO_DUTE.2.0"          <-----
    address 2.0.56.6/24
    secondary 52.0.56.6/24
    port 1/1/2:2
    mac 00:00:00:00:00:06
    ipv6
        address 3ffe::200:3806/120 primary-preference 1
        address 3ffe::3400:3806/120 primary-preference 2
    exit
    if-attribute                        <-----
        srlg-group "SrlgGroup_3"      <-----
    exit
    no shutdown
exit
interface "loopbackF.1.0"
    address 1.0.66.6/32
    secondary 51.0.66.6/32
    loopback
    ipv6
        address 3ffe::100:4206/128 primary-preference 1
        address 3ffe::3300:4206/128 primary-preference 2
    exit
    no shutdown
exit
interface "loopbackF.2.0"
    address 2.0.66.6/32
    secondary 52.0.66.6/32
    loopback
    ipv6
        address 3ffe::200:4206/128 primary-preference 1
        address 3ffe::3400:4206/128 primary-preference 2
    exit
    no shutdown
exit
interface "system"
    address 10.20.1.6/32
    ipv6
        address 3ffe::a14:106/128
    exit
    no shutdown
exit

```

```

        ip-fast-reroute
        router-id 10.20.1.6
#-----
echo "MPLS Label Range Configuration"
#-----
        mpls-labels
        sr-labels start 20000 end 80000
        exit
#-----
echo "OSPFv2 Configuration"
#-----
        ospf 0 10.20.1.6
        traffic-engineering
        database-export identifier 0
        advertise-router-capability area
        loopfree-alternates <-----
            remote-lfa <-----
                node-protect <-----
            exit <-----
            ti-lfa max-sr-frr-labels 3 <-----
                node-protect <-----
            exit <-----
        exit <-----
        segment-routing
        prefix-sid-range start-label 70000 max-index 999
        egress-statistics
            adj-set
            adj-sid
            node-sid
        exit
        ingress-statistics
            adj-set
            adj-sid
            node-sid
        exit
        no shutdown
    exit
    area 0.0.0.0
        interface "system"
            node-sid index 9
            no shutdown
        exit
        interface "DUTF_TO_DUTC.1.0" <-----
            interface-type point-to-point
            hello-interval 2
            dead-interval 10
            metric 10
            lfa-policy-map route-nh-template "templateSrlgGroup_1" <-----
            no shutdown
        exit
        interface "DUTF_TO_DUTD.1.0"
            interface-type point-to-point
            hello-interval 2
            dead-interval 10
            metric 1000
            lfa-policy-map route-nh-template "templateSrlgGroup_2"
            no shutdown
        exit
        interface "DUTF_TO_DUTE.1.0"
            interface-type point-to-point
            hello-interval 2
            dead-interval 10
            metric 10
            lfa-policy-map route-nh-template "templateSrlgGroup_3"

```

```

        no shutdown
    exit
    interface "loopbackF.1.0"
        node-sid index 3
        no shutdown
    exit
    interface "DUTF_TO_DUTC.2.0"
        interface-type point-to-point
        hello-interval 2
        dead-interval 10
        metric 10
        lfa-policy-map route-nh-template "templateSrlgGroup_4"
        no shutdown
    exit
    interface "DUTF_TO_DUTD.2.0"
        interface-type point-to-point
        hello-interval 2
        dead-interval 10
        metric 1000
        lfa-policy-map route-nh-template "templateSrlgGroup_5"
        no shutdown
    exit
    interface "DUTF_TO_DUTE.2.0"
        interface-type point-to-point
        hello-interval 2
        dead-interval 10
        metric 10
        lfa-policy-map route-nh-template "templateSrlgGroup_6"
        no shutdown
    exit
    interface "loopbackF.2.0"
        node-sid index 15
        no shutdown
    exit
    exit
    no shutdown
exit

```

*A:Dut-F# tools dump router segment-routing tunnel

=====

====

Legend: (B) - Backup Next-hop for Fast Re-Route

(D) - Duplicate

label stack is ordered from top-most to bottom-most

=====

====

-----+-----

Prefix	Sid-Type	Fwd-Type	In-Label	Prot-Inst	Out-Label(s)	Interface/
		Next Hop(s)				
Tunnel-ID						
1.0.33.3	Node	Orig/Transit	70000	OSPF-0	<-----	
1.0.36.3		1.0.36.3			<-----	
DUTC.1.0					40000	DUTF_TO_

	(B)1.0.56.5			60000	DUTF_T0_
DUTE.1.0 <-----					
1.0.44.4					
Node	Orig/Transit	70001	OSPF-0	<-----	
	1.0.36.3			40001	DUTF_T0_
DUTC.1.0 <-----					
	(B)1.0.46.4			50001	DUTF_T0_
DUTD.1.0 <-----					
1.0.55.5					
Node	Orig/Transit	70002	OSPF-0		
	1.0.56.5			60002	DUTF_T0_
DUTE.1.0					
	(B)1.0.36.3			40002	DUTF_T0_
DUTC.1.0					
1.0.66.6					
Node	Terminating	70003	OSPF-0		
1.0.11.1					
Node	Orig/Transit	70004	OSPF-0		
	1.0.36.3			40004	DUTF_T0_
DUTC.1.0					
	(B)1.0.46.4			50004	DUTF_T0_
DUTD.1.0					
1.0.22.2					
Node	Orig/Transit	70005	OSPF-0		
	1.0.36.3			40005	DUTF_T0_
DUTC.1.0					
	(B)1.0.46.4			50005	DUTF_T0_
DUTD.1.0					
10.20.1.3					
Node	Orig/Transit	70006	OSPF-0		
	1.0.36.3			40006	DUTF_T0_
DUTC.1.0					
	(B)1.0.56.5			60006	DUTF_T0_
DUTE.1.0					
10.20.1.4					
Node	Orig/Transit	70007	OSPF-0		
	1.0.36.3			40007	DUTF_T0_
DUTC.1.0					
	(B)1.0.46.4			50007	DUTF_T0_
DUTD.1.0					
10.20.1.5					
Node	Orig/Transit	70008	OSPF-0		
	1.0.56.5			60008	DUTF_T0_
DUTE.1.0					
	(B)1.0.36.3			40008	DUTF_T0_
DUTC.1.0					
10.20.1.6					
Node	Terminating	70009	OSPF-0		
10.20.1.1					
Node	Orig/Transit	70010	OSPF-0		
	1.0.36.3			40010	DUTF_T0_
DUTC.1.0					
	(B)1.0.46.4			50010	DUTF_T0_
DUTD.1.0					
10.20.1.2					
Node	Orig/Transit	70011	OSPF-0		
	1.0.36.3			40011	DUTF_T0_
DUTC.1.0					
	(B)1.0.46.4			50011	DUTF_T0_
DUTD.1.0					
2.0.33.3					
Node	Orig/Transit	70012	OSPF-0		
	1.0.36.3			40012	DUTF_T0_
DUTC.1.0					

DUTE.1.0	(B)1.0.56.5			60012	DUTF_T0_
2.0.44.4					
Node	Orig/Transit	70013	OSPF-0		
	1.0.36.3			40013	DUTF_T0_
DUTC.1.0	(B)1.0.46.4			50013	DUTF_T0_
DUTD.1.0					
2.0.55.5					
Node	Orig/Transit	70014	OSPF-0		
	1.0.56.5			60014	DUTF_T0_
DUTE.1.0	(B)1.0.36.3			40014	DUTF_T0_
DUTC.1.0					
2.0.66.6					
Node	Terminating	70015	OSPF-0		
2.0.11.1					
Node	Orig/Transit	70016	OSPF-0		
	1.0.36.3			40016	DUTF_T0_
DUTC.1.0	(B)1.0.46.4			50016	DUTF_T0_
DUTD.1.0					
2.0.22.2					
Node	Orig/Transit	70017	OSPF-0		
	1.0.36.3			40017	DUTF_T0_
DUTC.1.0	(B)1.0.46.4			50017	DUTF_T0_
DUTD.1.0					
2.0.56.5					
Adjacency	Transit	524282	OSPF-0		
	2.0.56.5			3	DUTF_T0_
DUTE.2.0	(B)1.0.56.5			3	DUTF_T0_
DUTE.1.0					
2.0.46.4					
Adjacency	Transit	524283	OSPF-0		
	2.0.46.4			3	DUTF_T0_
DUTD.2.0	(B)1.0.36.3			40001	DUTF_T0_
DUTC.1.0					
2.0.36.3					
Adjacency	Transit	524284	OSPF-0		
	2.0.36.3			3	DUTF_T0_
DUTC.2.0	(B)1.0.36.3			3	DUTF_T0_
DUTC.1.0					
1.0.56.5					
Adjacency	Transit	524285	OSPF-0		
	1.0.56.5			3	DUTF_T0_
DUTE.1.0	(B)1.0.36.3			40002	DUTF_T0_
DUTC.1.0					
1.0.46.4					
Adjacency	Transit	524286	OSPF-0		
	1.0.46.4			3	DUTF_T0_
DUTD.1.0	(B)1.0.36.3			40001	DUTF_T0_
DUTC.1.0					
1.0.36.3					
Adjacency	Transit	524287	OSPF-0		
	1.0.36.3			3	DUTF_T0_
DUTC.1.0 <----	(B)1.0.56.5			60000	DUTF_T0_
DUTE.1.0 <----					

```

-----+
No. of Entries: 24
-----+
*A:Dut-F#
*A:Dut-F# tools dump router ospf sr-database
=====
Rtr Base OSPFv2 Instance 0 Segment Routing Database
=====
SID      Label St Type Prefix          AdvRtr          Area  Flags          Stitching
-----+-----+-----+-----+-----+-----+-----+-----+
0         70000 +R   T1 1.0.33.3/32          10.20.1.3      0.0.0.0 [NnP]    ] T(R)    - <-----
1         70001 +R   T1 1.0.44.4/32          10.20.1.4      0.0.0.0 [NnP]    ] Tn(R)   - <-----
2         70002 +R   T1 1.0.55.5/32          10.20.1.5      0.0.0.0 [NnP]    ] L(R)    - <-----
3         70003 +R  LT1 1.0.66.6/32         10.20.1.6      0.0.0.0 [NnP]    ] -       -
4         70004 +R   T1 1.0.11.1/32         10.20.1.1      0.0.0.0 [NnP]    ] Tn(R)   -
5         70005 +R   T1 1.0.22.2/32          10.20.1.2      0.0.0.0 [NnP]    ] Tn(R)   -
6         70006 +R   T1 10.20.1.3/32         10.20.1.3      0.0.0.0 [NnP]    ] T(R)    -
7         70007 +R   T1 10.20.1.4/32         10.20.1.4      0.0.0.0 [NnP]    ] Tn(R)   -
8         70008 +R   T1 10.20.1.5/32         10.20.1.5      0.0.0.0 [NnP]    ] L(R)    -
9         70009 +R  LT1 10.20.1.6/32         10.20.1.6      0.0.0.0 [NnP]    ] -       -
10        70010 +R   T1 10.20.1.1/32         10.20.1.1      0.0.0.0 [NnP]    ] Tn(R)   -
11        70011 +R   T1 10.20.1.2/32         10.20.1.2      0.0.0.0 [NnP]    ] Tn(R)   -
12        70012 +R   T1 2.0.33.3/32          10.20.1.3      0.0.0.0 [NnP]    ] L(R)    -
13        70013 +R   T1 2.0.44.4/32          10.20.1.4      0.0.0.0 [NnP]    ] Tn(R)   -
14        70014 +R   T1 2.0.55.5/32          10.20.1.5      0.0.0.0 [NnP]    ] L(R)    -
15        70015 +R  LT1 2.0.66.6/32         10.20.1.6      0.0.0.0 [NnP]    ] -       -
16        70016 +R   T1 2.0.11.1/32         10.20.1.1      0.0.0.0 [NnP]    ] Tn(R)   -
17        70017 +R   T1 2.0.22.2/32          10.20.1.2      0.0.0.0 [NnP]    ] Tn(R)   -
-----+-----+-----+-----+-----+-----+
No. of Entries: 18
-----+
St:  R:reported I:incomplete W:wrong N:not reported F:failed
     +:SR-ack -:no route
Type: L:local M: mapping Srv Tx: route type
FRR:  L:Lfa R:RLfa T:TiLfa (R):Reported (F):Failed
      Ln, Rn, Tn: FRR providing node-protection
=====
*A:Dut-F#

```

2.1.9.5 LFA protection using a segment routing backup node SID

One of the challenges in MPLS deployments across multiple IGP areas or domains, such as in seamless MPLS design, is the provisioning of FRR local protection in access and metropolitan domains that make use of a ring, a square, or a partial mesh topology. To implement IP, LDP, or SR FRR in these topologies, the remote LFA feature must be implemented. Remote LFA provides a Segment Routing (SR) tunneled LFA next hop for an IP prefix, an LDP tunnel, or an SR tunnel. For prefixes outside of the area or domain, the access or aggregation router must push the following labels:

- service label
- BGP label for the destination PE
- LDP/RSVP/SR label to reach the exit ABR or ASBR
- label for the remote LFA next hop

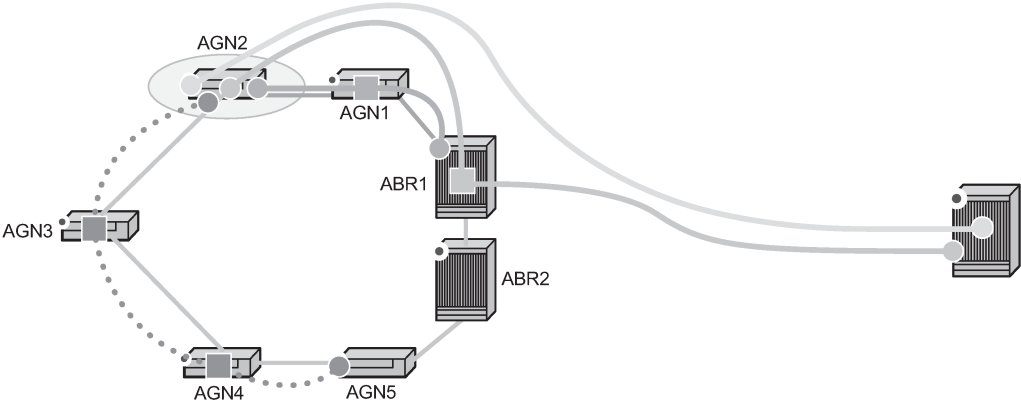
Small routers deployed in these parts of the network have limited MPLS label stack size support.

Figure 12: Label stack for remote LFA in ring topology shows the label stack required for the primary next hop and the remote LFA next hop computed by aggregation node AGN2 for the inter-area prefix of a remote PE. For an inter-area BGP label unicast route prefix for which ABR1 is the primary exit ABR, AGN2 resolves the prefix to the transport tunnel of ABR1 and, therefore, uses the remote LFA next hop of ABR1 for protection. The primary next hop uses two transport labels plus a service label. The remote LFA next hop for ABR1 uses PQ node AGN5 and pushes three transport labels plus a service label.

Seamless MPLS with Fast Restoration requires up to four labels to be pushed by AGN2, as shown in the following figure.

Figure 12: Label stack for remote LFA in ring topology

Label Location	Label Name	Assigned By	Protocol	Use
Label 1 (Bottom)	Service (PW, VC) Label	Remote PE	MP-BGP, T-LDP	Identifies Service on Remote PE
Label 2	Inter-Area Label	ABR1	BGP-LU	Identifies Path to Remote PE
Label 3	Intra-Area	AGN1	LDP, RSVP, SR	Identifies Path to ABR1
Label 4 (Top)	R-LFA Label	AGN3	LDP, RSVP, SR	Identifies Path to AGN5



0935

The objective of the LFA protection with a backup node SID feature is to reduce the label stack pushed by AGN2 for BGP label unicast inter-area prefixes. When link AGN2-AGN1 fails, packets are directed away from the failure and forwarded toward ABR2, which acts as the backup for ABR1 (and the other way around when ABR2 is the primary exit ABR for the BGP label unicast inter-area prefix). This requires that

ABR2 advertise a special label for the loopback of ABR1 that attracts packets normally destined for ABR1. These packets are forwarded by ABR2 to ABR1 via the inter-ABR link.

As a result, AGN2 pushes the label advertised by ABR2 to back up ABR1 on top of the BGP label for the remote PE and the service label. This keeps the label stack the same size for the LFA next hop to be the same size as that of the primary next hop. It is also the same size as the remote LFA next hop for the local prefix within the ring.

2.1.9.5.1 Configuring LFA using a backup node SID in OSPF

LFA using a backup node SID is enabled by configuring a backup node SID at an ABR/ASBR that acts as a backup to the primary exit ABR/ASBR of inter-area/inter-AS routes learned as BGP labeled routes.

```
config>router>ospf>segment-routing$  
  - backup-node-sid ip-prefix/prefix-length index 0..4294967295  
  - backup-node-sid ip-prefix/prefix-length label 1..4294967295
```

The user can enter either a label or an index for the backup node SID.



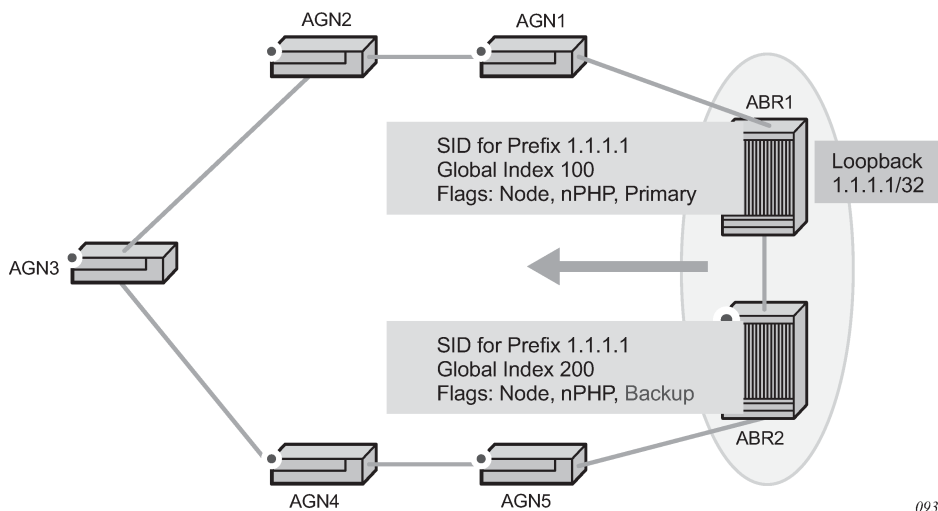
Note: This feature only allows the configuration of a single backup node SID per OSPF instance and per ABR/ASBR. In other words, only one pair of ABR/ASBR nodes can back up each other in an OSPF domain. Each time the user invokes the preceding command within the same OSPF instance, it overrides any previous configuration of the backup node SID. The same ABR/ASBR can, however, participate in multiple OSPF instances and provide a backup support within each instance.

2.1.9.5.2 Detailed operation of LFA protection using a backup node SID

As shown in the following figure, LFA for seamless MPLS supports environments where the boundary routers are either:

- ABR nodes that connect with Interior Border Gateway Protocol (IBGP) multiple domains, each using a different area of the same IGP instance
- ASBR nodes that connect domains running different IGP instances and use IBGP within a domain and External Border Gateway Protocol (EBGP) to the other domains

Figure 13: Backup ABR node SID



The following steps describe the configuration and behavior of LFA protection using a backup node SID.

1. The user configures node SID 100 in ABR1 for its loopback prefix 1.1.1.1/32. This is the regular node SID. ABR1 advertises the prefix SID sub-TLV for this node SID in the IGP and installs the ILM using a unique label.
2. Each router receiving the prefix sub-TLV for node SID 100 resolves it as described in [Segment routing in shortest path forwarding](#). Changes to the programming of the backup NHLFE of node SID 100, based on receiving the backup node SID for prefix 1.1.1.1/32, are defined in [Duplicate SID handling](#).
3. The user configures a backup node SID 200 in ABR2 for the loopback 1.1.1.1/32 of ABR1. The SID value must be different from that assigned by ABR1 for the same prefix. ABR2 installs the ILM, which performs a swap operation from the label of SID 200 to that of SID 100. The ILM must point to a direct link and next hop to reach 1.1.1.1/32 of ABR1 as its primary next hop. The IGP examines all adjacencies established in the same area as that of prefix 1.1.1.1/32 and determines which ones have ABR1 as a direct neighbor and with the best cost. If more than one adjacency has the best cost, the IGP selects the one with the lowest interface index. If there is no adjacency to reach ABR2, the prefix SID for the backup node is flushed and is not resolved, which prevents any other non-direct path being used to reach ABR1. As a result, any received traffic on the ILM of SID 200 traffic is blackholed.
4. If resolved, ABR2 advertises the prefix SID sub-TLV for backup node SID 200 and indicates in the SR Algorithm field that a modified SPF algorithm, referred to as "Backup-constrained-SPF", is required to resolve this node SID.
5. Each router receiving the prefix sub-TLV for backup node SID 200 and performs the following resolution steps.



Note: The following resolution steps do not require a CLI command to be enabled.

- a. The router determines which router is being backed up. This is achieved by checking the router ID owner of the prefix sub-TLV that was advertised with the same prefix but without the backup flag and which is used as the best route for the prefix. In this case, ABR1 is the router being backed up. Then the router runs a modified SPF by removing node ABR1 from the topology to resolve backup node SID 200. The primary next hop points to the path to ABR2 in the counterclockwise direction of the ring.

The router does not compute an LFA or a remote LFA for back node SID 200 because the main SPF used a modified topology.

- b. The router installs the ILM and primary NHLFE for the backup node SID.

Only a swap label operation is configured by all routers for the backup node SID. There is no push operation, and no tunnel for the backup node SID is added into the TTM.

- c. The router programs the backup node SID as the LFA backup for the SR tunnel to node SID of 1.1.1.1/32 of ABR1. In other words, each router overrides the remote LFA backup for prefix 1.1.1.1/32, which is normally PQ node AGN5.
 - d. If a router, such as AGN1, is adjacent to ABR1, it also programs the backup node SID as the LFA backup for the protection of any adjacency SID to ABR1.
6. When node AGN2 resolves a BGP labeled route for an inter-area prefix for which the primary ABR exit router is ABR1, it uses the backup node SID of ABR1 as the remote LFA backup instead of the SID to the PQ node (AGN5 in this example) to save on the pushed label stack.
- AGN2 continues to resolve the prefix SID for any remote PE prefix that is summarized into the local area of AGN2, as usual. AGN2 programs a primary next hop and a remote LFA next hop. Remote LFA uses AGN5 as the PQ node and pushes two labels, as it would for an intra-area prefix SID. There is no need to use the backup node SID for this prefix SID and force its backup path to go to ABR1. The backup path may exit from ABR2 if the cost from ABR2 to the destination prefix is shorter.
- 7. If the user excludes a link from LFA in the IGP instance using the **config>router>ospf>area>interface>loopfree-alternate-exclude** command, a backup node SID that resolves to that interface is not used as a remote LFA backup in the same way as regular LFA or PQ remote LFA next hop behavior.
 - 8. If the OSPF neighbor of a router is put into overload or if the metric of an OSPF interface to that neighbor is set to LSInfinity (0xFFFF), a backup node SID that resolves to that neighbor is not used as a remote LFA backup in the same way as regular LFA or PQ remote LFA next hop behavior.
 - 9. The LFA policy is supported with a backup node SID. See [Application of LFA policy to backup node SID tunnel](#).

2.1.9.5.3 Duplicate SID handling

When the IGP issues or receives an LSA or LSP containing a prefix SID sub-TLV for a node SID or a backup node SID with a SID value that is a duplicate of an existing SID or backup node SID, the resolution as described in the following table is followed.

Table 4: Handling of duplicate SIDs

	New LSA/LSP			
Old LSA/LSP	Backup node SID	Local backup node SID	Node SID	Local node SID
Backup Node SID	Old	New	New	New
Local Backup Node SID	Old	Equal	New	New

	New LSA/LSP			
Old LSA/LSP	Backup node SID	Local backup node SID	Node SID	Local node SID
Node SID	Old	Old	Equal/Old ¹	Equal/New ²
Local Node SID	Old	Old	Equal/Old ¹	Equal/Old ¹

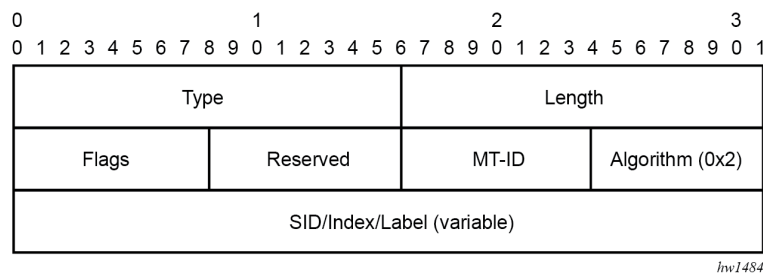
2.1.9.5.4 OSPF control plane extensions

All routers supporting OSPF control plane extensions must advertise support of the Backup-constrained-SPF algorithm of value 2 in the SR-Algorithm TLV, which is advertised in the Router Information Opaque LSA. This is in addition to the default supported algorithm "IGP-metric-based-SPF" of value 0. The following figure shows the encoding of the prefix SID sub-TLV to indicate a node SID of type backup and to indicate the modified SPF algorithm in the SR Algorithm field.



Note: The values used in the Flags field and in the Algorithm field are SR OS proprietary. The new Algorithm (0x2) field and values are used by this feature.

Figure 14: OSPF Prefix SID sub-TLV



The following table lists the OSPF Prefix SID sub-TLV main field values.

Table 5: OSPF Prefix SID sub-TLV main fields

Field	Value
Type	2
Length	variable

¹ Equal/Old means the following.

- If the prefix is duplicate, it is equal and no change is needed. Keep the old LSA/LSP.
- If the prefix is not duplicate, still keep the old LSA/LSP.

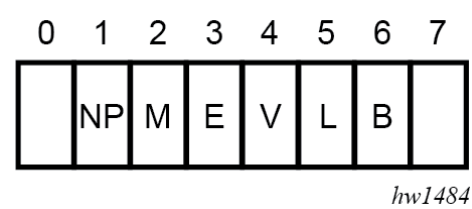
² Equal/New means the following.

- If the prefix is duplicate, it is equal and no change is needed. Keep the old LSA/LSP.
- If the prefix is not duplicate, pick a new prefix and use the new LSA/LSP.

Field	Value
Flags	1 octet field

The following figure shows the details of the OSPF Prefix SID sub-TLV flag field; the B-flag is new and SR OS proprietary.

Figure 15: OSPF Prefix SID sub-TLV flags



The following table describes OSPF Prefix SID sub-TLV flags, including the new B-flag.

Table 6: OSPF Prefix SID sub-TLV flags

Flag	Description
NP-Flag	No-PHP flag If set, the penultimate hop must not pop the prefix SID before delivering the packet to the node that advertised the prefix SID.
M-Flag	Mapping Server Flag If set, the SID is advertised from the segment routing mapping server functionality as described in <i>draft-filsfils-spring-segment-routing-ldp-interop</i> .
E-Flag	Explicit-Null Flag If set, any upstream neighbor of the prefix SID originator must replace the prefix SID with a prefix SID having an Explicit-NULL value (0 for IPv4) before forwarding the packet.
V-Flag	Value/Index Flag If set, the prefix SID carries an absolute value. If not set, the prefix SID carries an index.
L-Flag	Local/Global Flag If set, the value or index carried by the prefix SID has local significance. If not set, then the value or index carried by this sub-TLV has global significance.
B-Flag	This flag is used by the Protection using backup node SID feature. If set, the SID is a backup SID for the prefix. This value is SR OS proprietary.
Other bits	Reserved

Flag	Description
	These must be zero when sent and are ignored when received.
MT-ID	Multitopology ID, as defined in RFC 4915.
Algorithm	This octet identifies the algorithm that the prefix SID is associated with. A value of (0x2) indicates the modified SPF algorithm, which removes from the topology the node that is backed up by the backup node SID. This value is SR OS proprietary.
SID/Index/Label	Based on the V and L flags, this field contains either: - a 32-bit index defining the offset in the SID or Label space advertised by this router - a 24-bit label where the 20 rightmost bits are used for encoding the label value

2.1.9.6 Multihomed prefix LFA extensions in SR-OSPF

This feature makes use of the Multihomed Prefix (MHP) model described in RFC 8518 to compute a backup IP next-hop using an alternate ABR or ASBR for external prefixes and to an alternate router owner for local anycast prefixes.

The feature applies to OSPF routes of external /32 prefixes (OSPFv2 routes types 3, 4, 5, and 7) and local /32 anycast prefixes if the prefix is not protected by base LFA.

The computed IP next-hop based backup path is programmed for SR-OSPF node SID tunnels of external /32 prefixes and to /32 prefixes in same area as the computing node and which are advertised by multiple routers (anycast prefixes) in both algorithm 0 and flexible-algorithm numbers.

See “Multihomed prefix LFA extensions in OSPF” in the *7705 SAR Gen 2 Unicast Routing Protocols Guide* for more information about the configuration of this feature.

2.1.9.7 Multihomed prefix LFA extensions in SR IS-IS

This feature makes use of the Multihomed Prefix (MHP) model described in RFC 8518 to compute a backup IP next-hop using an alternate ABR or ASBR for external prefixes and to an alternate router owner for local anycast prefixes.

The algorithm described in RFC 8518 is limited in scope to only computed backup paths consisting of direct IP next hops and tunneled next hops (IGP shortcuts).

The computed backup paths are added to IS-IS routes of external /32 and /128 prefixes and intra-area /32 and /128 anycast prefixes in the Routing Table Manager (RTM) if the prefix is not protected by the base LFA.

The computed backup path is also programmed for the following tunnels:

- SR IS-IS IPv4 and IPv6 node SID tunnels of external /32 and /128 prefixes and of intra-area /32 and /128 anycast prefixes, in both algorithm 0 and flexible algorithm numbers

As a result, an SR-TE LSP or an SR-MPLS policy that uses an SR IS-IS SID of those same prefixes in its configured or computed SID list benefits from the multihomed prefix LFA protection.

See “Multihomed prefix LFA extensions in IS-IS” in the *7705 SAR Gen 2 Unicast Routing Protocols Guide* for more information about the configuration of this feature.

2.1.9.8 LFA solution across IGP area or instance boundary using SR repair tunnel in SR-OSPF

This feature enhances the IP next-hop based MHP backup path calculation specified in RFC 8518 with the addition of the support of an SR repair tunnel. The SR repair tunnel uses a PQ node or a P-Q set to reach the alternate exit ABR or ASBR for external prefixes, or alternate owner router for intra-area anycast prefixes. This capability is in addition to supporting the RFC 8518 algorithm used in the case where the path to prefix P using the alternate exit ABR or ASBR (or alternate owner router) is in the shortest path from the neighbor of the computing node.

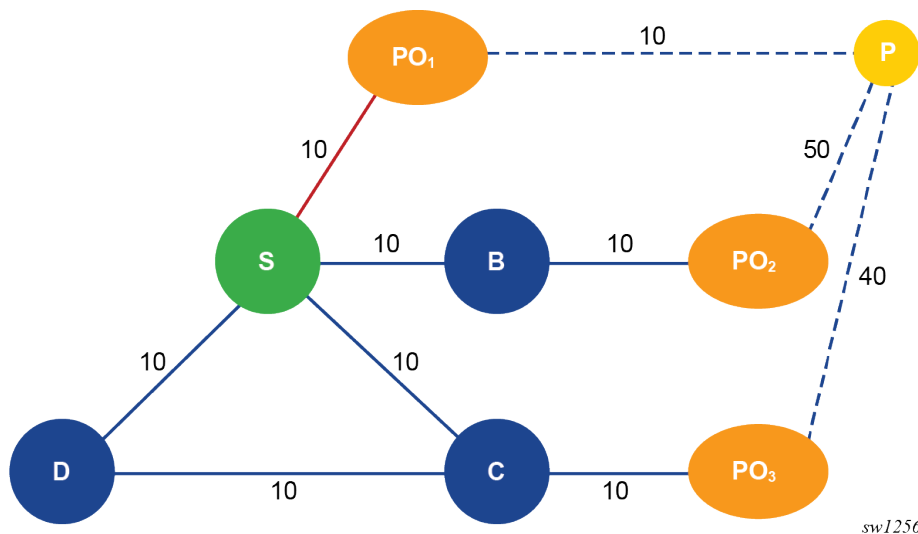
This feature applies the computed backup path to SR-OSPF node SID tunnels of external /32 prefixes and to /32 prefixes in the same area as the computing node, and which are advertised by multiple routers (anycast prefixes) in both algorithm 0 and flexible-algorithm numbers. It also extends the protection to any SR-TE LSP or SR policy that uses an SR-OSPF SID of those same prefixes in its configured or computed SID list.

This feature shares the same configuration CLI commands as the MHP LFA feature as described in [Multihomed prefix LFA extensions in SR-OSPF](#). After the IP next-hop based MHP LFA is enabled, the extensions to compute an SR repair tunnel for the MHP LFA in the case of SR-OSPF are automatically enabled if the user enables TI-LFA or RLFA. The computation reuses the SID list of the primary path or of the TI-LFA or RLFA backup path of the alternate ABR, ASBR, or alternate owner router. The algorithm details are described in [Extending MHP LFA coverage with repair tunnels for SR OSPF](#).

2.1.9.8.1 Extending MHP LFA coverage with repair tunnels for SR OSPF

The following figure shows topology that is used as a reference in this section.

Figure 16: Application of MHP LFA to SR-OSPF tunnel of external prefix



For computing node S, PO_1 is the ABR in the best path (PO_{best}) to reach prefix P. None of the neighbors of node S satisfies the link or node protection inequality of RFC 8518 described in "RFC 8518 multihomed prefix LFA for OSPF" in the *7705 SAR Gen 2 Unicast Routing Protocols Guide*. Therefore, the main aspect of the extension to the RFC 8518 algorithm is for node S to find the best repair tunnel using a PQ node or a P-Q set, which forwards the packet to an alternate exit ABR or ASBR represented by node PO_1 , PO_2 , or PO_3 in [Figure 16: Application of MHP LFA to SR-OSPF tunnel of external prefix](#).



Note: The same calculation is applied to intra-area /32 anycast prefixes and in that case POi nodes represent the multiple owner routers of the prefix.

The following are the steps of this algorithm.

1. Compute a multihomed LFA repair tunnel for prefix P using each POi.

a. Node S first attempts to compute a MHP LFA repair tunnel path that matches one ECMP primary path to a POi and that avoids neighbor node E. In other words, the repair tunnel uses POi as a PQ node. Node S further restricts the set of ECMP paths to those over an outgoing interface that satisfy any LFA policy applied to link S-E. Specifically Node S:

- excludes paths that do not satisfy the admin-group or SRLG constraint in the LFA policy of the primary next hop to E of prefix P
- applies preference of IP next hops versus tunneled next hops (IGP shortcuts), in accordance with the configuration of the LFA policy and prefers tunneled next hops terminating on the POi node, regardless the protection level
- prefers the LFA next hop not sharing the same pseudo-node (PN) as the primary next hop
- applies preference of node protection versus link protection as per the configuration of the LFA policy
- applies the admin-group preference configured in the LFA policy
- selects the next hops with the lowest IGP cost to the destination prefix P
- selects the tunnel closest (lowest IGP cost) to the destination among equal cost tunnel next hops
- selects the LFA neighbor with the lowest router ID among equal cost tunneled or IP next hops
- selects the lowest tunnel ID or interface ID among next hops to the same LFA neighbor

See [LFA solution across IGP area or instance boundary using SR repair tunnel in SR-OSPF](#) for more information about the algorithm interaction with the LFA policy feature.

b. If no path is found in step (1.a), node S computes a MHP LFA repair-tunnel path that matches the node-protect or link-protect LFA, TI-LFA, or RLFA backup path of node POi. In this case, the MHP LFA repair tunnel effectively uses a PQ node or a P-Q set to force the packet to exit the local area at the selected POi.



Note: If all ECMP candidate paths in step (1.a) are excluded by applying the LFA policy of link S-E, no LFA, TI-LFA, or RLFA backup path of node POi is found in this step because ECMP and LFA are mutually exclusive per prefix.

c. If no candidate path is found in steps (1.a) and (1.b), POi is not a candidate alternate ABR, alternate ASBR, or alternate owner router.

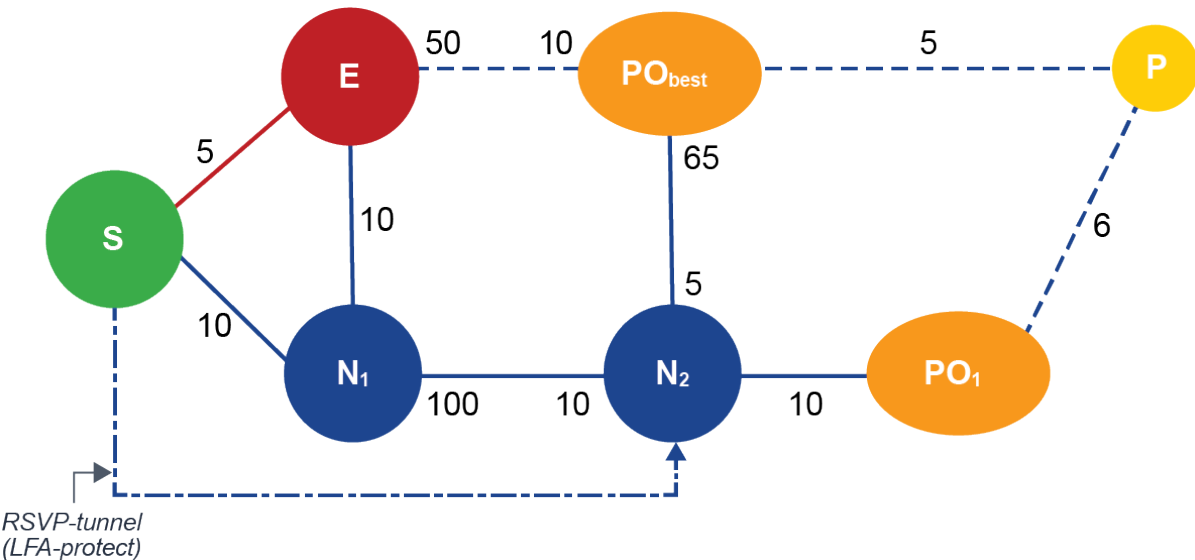
2. Create an ordered list of candidate MHP LFA tunnel paths with the following preference order (from highest to lowest).

- a. Prefer the candidate path that uses a POi with the next hop of the primary path, avoiding neighbor node E. Candidate paths are split into two subsets and paths computed from step 1.a preferred over paths computed from step 1.b.
 - b. Within each subset, prefer the candidate path that uses POi with lower total cost to prefix P expressed as $\text{Min}\{D_{\text{opt}}(S, \text{POi}) + \text{cost}(\text{POi}, P)\}$.
 - c. If the cost is the same, prefer the candidate path that uses a POi with the lower label stack size.
 - d. If the label stack size is the same, prefer the candidate path that uses a POi with the lower router ID.
3. Analyze the ordered list and select the first MHP LFA tunnel path with a segment list size that does not exceed either the value of 1 if RLFA is enabled but TI-LFA is disabled, or the value of the **loopfree-alternate ti-lfa max-sr-frr-labels** command if TI-LFA is enabled.
4. Program in datapath the segment list of the selected MHP LFA repair tunnel for the specific prefix P. The segment list consists of pushing on top of the SID of destination prefix P the SID of the PQ node or the SIDs of the P-Q set.

2.1.9.8.2 Example application of MHP LFA with repair tunnel

The following figure shows topology that is used as a reference in this section.

Figure 17: Application of MHP LFA with repair tunnel to SR-OSPF tunnel of external or anycast prefix



Prefix	Sid-Type	Fwd-Type	Next Hop (s)	Out-Label (s)	Interface/Tunnel-ID
10.20.1.1	Node	Terminating			
10.20.1.2	Node	Orig/Transit	1.1.2.2	19020	to_Dut-B
		(B) 1.1.3.3	19020		to_Dut-C
10.20.1.3	Node	Orig/Transit	1.1.3.3	19030	to_Dut-C
		(B) 1.1.2.2	19030		to_Dut-B
10.20.1.4	Node	Orig/Transit	10.20.1.4	19040	1 (RSVP)
10.20.1.5	Node	Orig/Transit	1.1.3.3	19050	to_Dut-C
		(B) 1.1.2.2	19050		to_Dut-B
10.20.1.6	Node	Orig/Transit	10.20.1.4	19060	1 (RSVP)
6.6.6.6	Node	Orig/Transit	1.1.3.3	19160	to_Dut-C
		(B) 10.20.1.4	19060		1 (RSVP)
			19160		

[S] dut-A node SID 10
[N₁] dut-B node SID 20
[E] dut-C node SID 30
[N₂] dut-D node SID 40
[PO_{best}] dut-E node SID 20
[PO₁] dut-F node SID 60
P] 6.6.6.6/32 anycast SID 160

sw1257

Node S is connected to nodes E and N₁ using IP links, and to node N₂ using an IGP shortcut (RSVP-TE LSP).

Prefix P (6.6.6.6/32) is either:

- an external prefix with prefix SID re-advertised by ASBR nodes PO_{best} and PO₁ and with best path through PO_{best}
or
- an anycast prefix with prefix SID owned by both routers PO_{best} and PO₁ with best path from node S is to PO_{best}

An SRGB assigned to the OSPF instance uses an offset label value of 19000. Base LFA, RLFA, TI-LFA, and MHP LFA are all enabled in node S. Node protection is also enabled. MHP LFA is preferred. Therefore, the following commands are enabled:

- **MD-CLI**

```
configure router ospf loopfree-alternate remote-lfa node-protect
configure router ospf loopfree-alternate ti-lfa node-protect
configure router ospf loopfree-alternate multi-homed-prefix preference all
```

- **classic CLI**

```
configure router ospf loopfree-alternates remote-lfa node-protect
configure router ospf loopfree-alternates ti-lfa node-protect
configure router ospf loopfree-alternates multi-homed-prefix preference all
```

The resulting LFA computations in node S for prefix P yield the following backup paths:

- base LFA node-protecting path to PO_{best} and using IGP shortcut to neighbor N_2 as next hop
- RLFA node-protecting path to PO_{best} and transiting through PQ node N_2
- TI-LFA node-protecting path to PO_{best} and transiting through PQ node N_2
- a MHP LFA path using the RFC 8518 node-protecting inequality as described in "RFC 8518 multihomed prefix LFA for OSPF" in the *7705 SAR Gen 2 Unicast Routing Protocols Guide*. This yields the same path as the base LFA, meaning a node-protecting path to PO_{best} and using IGP shortcut to neighbor N_2 as the next hop.

Node S does, however, determine that this path does not satisfy the PO_{best} overlap inequality as described in "Enhancement to RFC 8518 Algorithm for backup path overlap with path to PO_{best} in the local area" in the *7705 SAR Gen 2 Unicast Routing Protocols Guide* and, therefore, attempts an SR repair tunnel computation as the next step.

- MHP LFA path to PO_1 and using IGP shortcut to neighbor N_2 as next hop. This backup path forces the packet to arrive and exit (if P is external prefix) PO_1 by pushing PO_1 node SID with an index value of 60 and label value of 19060.

The MHP LFA repair tunnel is therefore the preferred backup path and is programmed in datapath to protect the primary path of prefix P.

2.1.9.9 LFA solution across IGP area or instance boundary using SR repair tunnel in SR IS-IS

This feature enhances the backup path calculation for the IP next-hop based multihomed path prefix in RFC 8518 with the addition of repair tunnels that make use of a PQ node or a P-Q set to reach the alternate exit ABR or ASBR of external prefixes or the alternate owner router for intra-area anycast prefixes.

The feature programs the computed backup path for the following tunnels:

- SR IS-IS node SID tunnels of external /32 IPv4 prefixes and /128 IPv6 prefixes, and node SID tunnels of intra-area /32 IPv4 anycast prefixes and /128 anycast IPv6 prefixes, in both algorithm 0 and flexible-algorithms

As a result, an SR-TE LSP or an SR-MPLS policy that uses an SR IS-IS SID of those same prefixes in its configured or computed SID list benefits from the multihomed prefix LFA protection.

After the IP next-hop based multihomed prefix LFA is enabled, the extensions to compute a SR-TE repair tunnel for the multihomed prefix LFA in the case of SR IS-IS are automatically enabled if the user also enabled TI-LFA or Remote LFA. The computation reuses the SID list of the primary path or of the TI-LFA or Remote LFA backup path of the alternate ABR or ASBR or alternate owner router.

The behavior of this feature is the same as in OSPF. See [LFA solution across IGP area or instance boundary using SR repair tunnel in SR-OSPF](#).

2.1.10 Segment routing datapath support

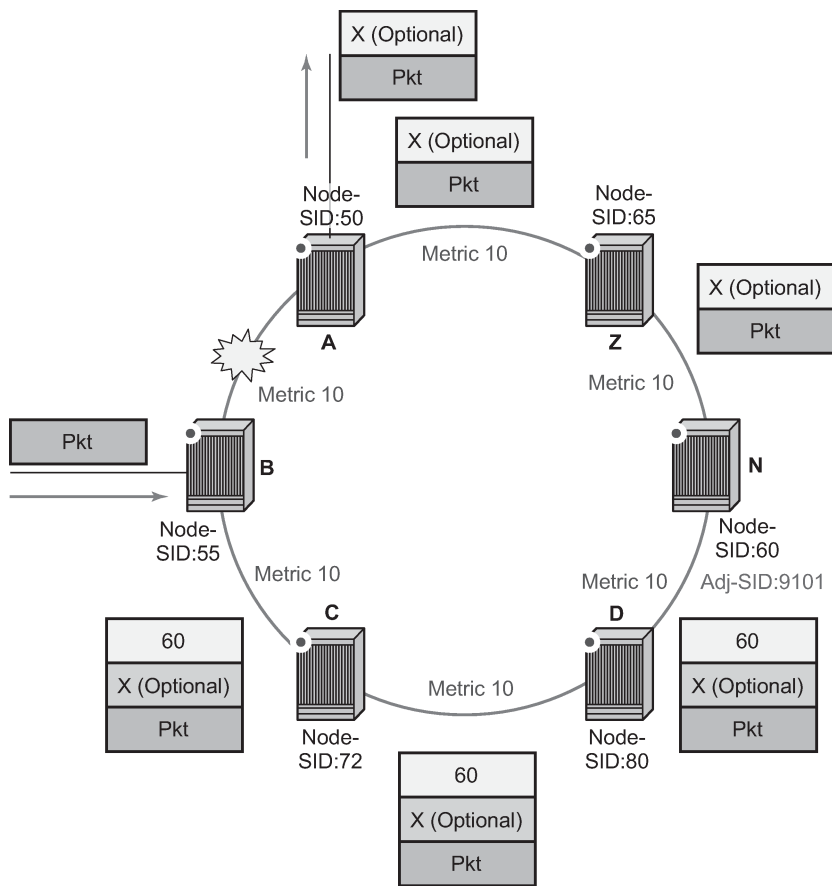
A packet received with a label matching either a node SID or an adjacency SID is forwarded according to the ILM type and operation, as described in the following table.

Table 7: Datapath support

Label type	Operation
Top label is a local node SID	Label is popped and the packet is further processed. If the popped node SID label is the bottom-of-stack label, the IP packet is looked up and forwarded in the appropriate FIB.
Top or next label is a remote node SID	Label is swapped to the calculated label value for the next hop and forwarded according to the primary or backup NHLFE. With ECMP, a maximum of 32 primary next-hops (NHLFEs) are programmed for the same destination prefix and for each IGP instance. ECMP and LFA next-hops are mutually exclusive, as in the current implementation.
Top or next label is an adjacency SID	Label is popped and the packet is forwarded from the interface to the next-hop associated with this adjacency SID label. This datapath operation is modeled like a swap to an implicit-null label instead of a pop.
Next label is BGP 3107 label	The packet is further processed according to the ILM operation, as in the current implementation. • The BGP label may be popped and the packet looked up in the appropriate FIB. • The BGP label may be swapped to another BGP label. • The BGP label may be stitched to an LDP label.
Next label is a service label	The packet is looked up and forwarded in the Layer 2 or VPRN FIB, as in the current implementation.

A router forwarding an IP or a service packet over a segment routing tunnel pushes a maximum of two transport labels with a remote LFA next hop, as shown in the following figure.

Figure 18: Transport label stack in shortest path forwarding with segment routing



al_0648

Assume that a VPRN service in node B forwards a packet received on a SAP to a destination VPN-IPv4 prefix X advertised by a remote PE2 via ABR/ASBR node A. Router B is in a segment routing domain while PE2 is in an LDP domain. BGP labeled routes are used to distribute the PE /32 loopbacks between the two domains.

When node B forwards over the primary next hop for prefix X, it pushes the node SID of the ASBR followed by the BGP 8277 label for PE2, followed by the service label for prefix X. When the remote LFA next hop is activated, node B pushes one or more segment routing label: the node SID for the remote LFA backup node (node N).

When node N receives the packet while the remote LFA next hop is activated, it pops the top segment routing label that corresponds to a local node SID. This results in popping this label and forwarding of the packet to the ASBR node over the shortest path (link N-Z).

When the ABR/ASBR node receives the packet from either node B or node Z, it pops the segment routing label that corresponds to a local node SID, then swaps the BGP label and pushes the LDP label of PE2, which is the next hop of the BGP labeled route.

2.1.10.1 Hash label and EL support

When the **hash-label** option is enabled in a service context, the hash label is always inserted at the bottom of the stack, in accordance with RFC 6391.

The LSR adds the capability to check a maximum of 16 labels in a stack. The LSR is able to hash on the IP headers when the payload below the label stack of maximum size of 16 is IPv4 or IPv6, including when a MAC header precedes it (**eth-encap-ip** command option).

The EL feature, as specified in RFC 6790, is supported on RSVP, LDP, segment-routed, and BGP transport tunnels. It uses the Entropy Label Indicator (ELI) to indicate the presence of the EL in the label stack. The ELI, followed by the actual EL, is inserted immediately below the transport label for which the EL feature is enabled. If multiple transport tunnels have the EL feature enabled, the ELI/EL is inserted below the lowest transport label in the stack.

The LSR hashing operates as follows:

- If the **lbl-only** hashing command option is enabled, or if one of the other LSR hashing options is enabled but a IPv4 or IPv6 header is not detected below the bottom of the label stack, the LSR hashes on the EL only.
- If the **lbl-ip** command option is enabled, the LSR hashes on the EL and the IP headers.
- If the **ip-only** or **eth-encap-ip** command option is enabled, the LSR hashes on the IP headers only.

For more information about the hash label and EL features, see the "MPLS EL and hash label" section in the *7705 SAR Gen 2 MPLS Guide*.

2.1.10.2 TTL or hop-limit field handling

The user can configure the TTL or hop-limit propagation for all Segment Routing MPLS (SR-MPLS) tunnels carrying IPv4 or IPv6 packets using the following CLI commands:

```
configure router ttl-propagate sr-mpls-local
configure router ttl-propagate sr-mpls-transit
```

This applies to IPv4 and IPv6 packets of IGP, BGP unlabeled (except 6PE), and static routes in the Base router whose next hop is resolved to an SR-MPLS tunnel of any of the following types:

- SR-ISIS
- SR-OSPF
- SR-OSPF3
- SR-TE
- LSP
- SR policy

By default handling, the IP TTL or hop limit is propagated to all labels in the segment routing transport label stack.

The user can configure the TTL or hop-limit propagation separately for CPM-originated IP packets and for transit IP packets. Transit IP packets are packets of Base router prefixes received on an access or network interface (with or without tunnel encapsulation), and whose lookup in the FIB results in forwarding them over an SR-MPLS tunnel.

More information about configuring the TTL or hop-limit propagation in other service or routing contexts is available as follows:

- See "Configuration of TTL propagation for VPRN routes" in the *7705 SAR Gen 2 Layer 3 Services Guide: IES and VPRN* for information about IPv4 and IPv6 packets forwarded in a VPRN or Layer 3 EVPN service context and resolved to an MPLS tunnel, including an SR-MPLS tunnel.
- See "Configuration of TTL propagation for BGP labeled routes" in the *7705 SAR Gen 2 Unicast Routing Protocols Guide* for information about IPv4 and IPv6 packets of routes in the Base router resolved to a BGP-LU or a 6PE tunnel, which itself resolves to an MPLS tunnel, including an SR-MPLS tunnel.

2.1.11 BGP shortcuts using segment routing tunnels

The user enables the resolution of IPv4 prefixes using SR tunnels to BGP next hops in the TTM by configuring the following command:

```
config>router>bgp>next-hop-resolution
  - shortcut-tunnel
    - [no] family {ipv4}
      - resolution {any | disabled | filter}
      - resolution-filter
        - [no] sr-isis
        - [no] sr-ospf
      - [no] disallow-igp
    - exit
  - exit
- exit
```

When **resolution** is set to **any**, any supported tunnel type in the BGP shortcut context is selected according to the TTM preference. The following tunnel types are supported in a BGP shortcut context in order of preference: RSVP, LDP, segment routing, and BGP.

When the **sr-isis** or **sr-ospf** command is enabled, an SR tunnel to the BGP next hop is selected in the TTM from the lowest preference IS-IS or OSPF instance. If many instances have the same lowest preference, the selection of the SR tunnel to the BGP next hop favors the lowest numbered IS-IS or OSPF instance.

See "BGP" in the *7705 SAR Gen 2 Unicast Routing Protocols Guide* for more information.

2.1.12 BGP labeled route resolution using segment routing tunnels

The user enables the resolution of RFC 8277 BGP labeled route prefixes using SR tunnels to BGP next hops in the TTM by configuring the following commands:

```
config>router>bgp>next-hop-resolution
  - labeled-routes
    - transport-tunnel
      - [no] family {label-ipv4 | label-ipv6 | vpn}
        - resolution {any | disabled | filter}
        - resolution-filter
          - [no] sr-isis
          - [no] sr-ospf
        - exit
      - exit
    - exit
  - exit
- exit
```

If **resolution** is set to **disabled**, the default binding to LDP tunnels is used. If **resolution** is set to **any**, any supported tunnel type in the BGP labeled route context is selected according to the TTM preference.

The following tunnel types are supported in a BGP labeled route context and are listed in order of preference:

1. RSVP
2. LDP
3. segment routing

When either **sr-isis** or **sr-ospf** is specified using the **resolution-filter** option, a tunnel to the BGP next hop is selected in the TTM from the lowest numbered IS-IS or OSPF instance.

See "BGP" in the *7705 SAR Gen 2 Unicast Routing Protocols Guide* for more information.

2.1.13 Service packet forwarding with segment routing

Two SDP subtypes of the MPLS type allow service binding to a segment routing tunnel programmed in the TTM by IS-IS or OSPF:

- **config>service>sdp>sr-isis**
- **config>service>sdp>sr-ospf**

An SDP configured as **sr-isis** or **sr-ospf** can be used with the **far-end** option. When the **sr-isis** or **sr-ospf** command is enabled, a tunnel to the far-end address is selected in the TTM from the lowest preference IS-IS or OSPF instance. The SR IS-IS or SR-OSPF tunnel is selected at the time of the binding, according to the tunnel selection rules. If a more preferred tunnel is subsequently added to the TTM, the SDP does not automatically switch to the new tunnel until the next time the SDP is being re-resolved.

The **tunnel-far-end** option is not supported. In addition, the **mixed-lsp-mode** option does not support the **sr-isis** and **sr-ospf** tunnel types.

The signaling protocol for the service labels for an SDP using a segment routing tunnel can be configured to static (**off**), T-LDP (**tl dp**), or BGP (**bgp**).

SR tunnels can be used in VPRN and BGP EVPN with the **auto-bind-tunnel** command. See "Next-hop resolution" in the *7705 SAR Gen 2 Unicast Routing Protocols Guide* for more information.

Both VPN-IPv4 and VPN-IPv6 (6VPE) are supported in a VPRN or BGP EVPN service using segment routing transport tunnels with the **auto-bind-tunnel** command.

See "BGP" in the *7705 SAR Gen 2 Unicast Routing Protocols Guide* and the *7705 SAR Gen 2 Layer 3 Services Guide: IES and VPRN* for more information about the VPRN **auto-bind-tunnel** CLI command.

2.1.14 Mirror services

The user can configure a spoke-SDP bound to an SR tunnel to forward mirrored packets from a mirror source to a remote mirror destination. In the configuration of the mirror destination service at the destination node, the **remote-source** command must use a spoke-sdp with VC-ID which matches the one the user configured in the mirror destination service at the mirror source node. The far-end option is not supported with an SR tunnel.

Configuration at mirror source node:

```
config mirror mirror-dest 10
```

```

- no spoke-sdp sdp-id:vc-id
- spoke-sdp sdp-id:vc-id [create]
  - egress
    - vc-label egress-vc-label

```

**Note:**

- *sdp-id* matches an SDP which uses an SR tunnel
- for *vc-label*, both static and t-ldp egress vc labels are supported

Configuration at mirror destination node:

```

*A:node:2# configure mirror mirror-dest 10 remote-source
- spoke-sdp <SDP-ID>:<VC-ID> create <-- VC-ID matching that of spoke-sdp configured in
  mirror destination context at mirror source node.
  - ingress
    - vc-label <ingress-vc-label> <--- optional: both static and t-ldp ingress vc
  label are supported.
  - exit
  - no shutdown
- exit
- exit

```

**Note:**

- the **far-end** command is not supported with SR tunnel at mirror destination node; user must reference a spoke-SDP using a segment routing SDP coming from mirror source node:
 - **far-end** *ip-address* [*vc-id vc-id*] [**ing-svc-label** *ingress-vc-label* | **tldp**] [**icb**]
 - **no far-end** *ip-address*
- for *vc-label*, both static and t-ldp ingress vc labels are supported

Mirroring is also supported with the PW redundancy feature when the endpoint spoke-sdp, including the ICB, is using an SR tunnel.

2.1.15 Class-based forwarding for SR-ISIS over RSVP-TE LSPs

To enable CBF+ECMP for SR-ISIS over RSVP-TE:

- Configure the resolution of SR over RSVP-TE LSPs as IGP shortcuts.
- Configure class based forwarding parameters in the MPLS context (a class forwarding policy, forwarding classes to sets associations, and RSVP-TE LSPs to forwarding sets associations).
- Enable class forwarding in the segment routing context.

When SR-ISIS resolves to an ECMP set of RSVP-TE LSPs and class forwarding is enabled in the segment routing context, the following behaviors apply:

- If no LSP in the full ECMP set, has been assigned with a class forwarding policy configuration, the set is considered as inconsistent from a CBF perspective. The system programs, in the forwarding path, the whole ECMP set and regular ECMP spraying occurs over the full set.
- If the ECMP set refers to more than one class forwarding policy, the set is inconsistent from a CBF perspective. The system programs, in the forwarding path, the whole ECMP set without any CBF information, and regular ECMP spraying occurs over the full set.

- In all other cases the ECMP set is considered consistent from a CBF perspective and the following rules apply:
 - If there is no default set (either user-defined or implicit) referenced in a CBF-consistent ECMP set, the system automatically selects one set as the default one. The selected set is the non-empty one with the lowest ID amongst those referenced by the LSPs of the ECMP set.
 - The system programs the data-path such that a packet which has been classified to a particular forwarding class is forwarded using the LSPs associated with the forwarding set which itself is associated with that forwarding class. In the event where the forwarding set is composed of multiple LSPs, the system performs ECMP over these LSPs.
 - Forwarding classes which are either not explicitly mapped to a set or which are mapped to a set for which all LSPs are down are forwarded using the default-set. The system re-elects a default set in cases where all the LSPs of the current default-set become inactive. The system also adapts (updates data-path programming) to configuration or state changes.
 - The CBF capability is available with any system profile. The number of sets is limited to four with system profile None or A, and to six with system profile B.

2.1.16 Segment routing traffic statistics

This section describes capabilities and procedures applicable to IS-IS, OSPFv2, and OSPFv3.

SR OS can enable and collect SID traffic statistics on the ingress and egress datapaths. Statistics can also be shown, monitored, and cleared, as well as accessed using telemetry.

IS-IS and OSPFv2 support Node SID, Adjacency SID, and Adjacency Set statistics. OSPFv3 supports Node SID and Adjacency SID statistics. The following commands are used to enter the context that allows for configuring the types of SIDs for which to collect traffic statistics:

- **configure router isis segment-routing egress-statistics**
- **configure router ospf segment-routing egress-statistics**
- **configure router ospf3 segment-routing egress-statistics**
- **configure router isis segment-routing ingress-statistics**
- **configure router ospf segment-routing ingress-statistics**
- **configure router ospf3 segment-routing ingress-statistics**

By default, statistics collection is disabled on all types of SIDs. If statistics are disabled after having been enabled, the statistics indexes that were allocated are released and the counter values are cleared.

On ingress, depending on which types of SIDs have statistics enabled, the system allocates a statistic index to each programmed ILM, corresponding to the following:

- the local node SID (including backup node SID) and the local adjacency SIDs (including adjacencies advertised as set members)
- the received node SID advertisements

On egress, depending on which types of SIDs have statistics enabled, the following apply:

- The system allocates a statistic index shared by the programmed NHLFEs (primary, and backup if any) corresponding to the local Adjacency SIDs and to the received Adjacency SIDs advertisements, and a statistic index shared by the primary NHLFEs (as many as members) of each adjacency set.

- The system allocates a statistic index shared by the programmed NHLFEs (one or more primaries, and backup if any) corresponding to each of the received node SID advertisements.



Note: The statistic indexes constitute a finite resource. The system may not be able to allocate as many indexes as needed. In this case, the system issues a notification and automatically retries to allocate statistic indexes, but does not issue further notifications in case it still fails to allocate the needed statistic indexes. If the system successfully allocates all the required statistic indexes to IGP SIDs, then a second notification is issued to inform the user. A state variable records whether a SID has an index allocated.



Note: The allocation of statistic indexes is non-deterministic. If more statistic indexes are required system-wide, for example, upon a reboot, the system may not be able to re-allocate the statistic indexes to the same entities as before the reboot.

2.1.17 Configuring BGP-based services for flexible algorithms

BGP-based network services (VPRN, EVPN and VPLS) can be automatically steered to a flexible algorithm using service-based import policies. To configure BGP automated steering, a policy-statement must first be defined. Within this policy-statement, all BGP criteria is identified to steer traffic towards certain prefixes towards a flexible algorithm topology. Often the BGP color community is used to identify which flex-algorithm is used for a BGP prefix, however, Nokia SR OS has the capability to match upon any existing policy-statement BGP attribute criteria.

The following example shows policy-statement configuration.

Example: MD-CLI

```
[ex: /configure router policy-options]
A:admin@node-2# info
policy-options {
    policy-statement "ExamplePolicy" {
        entry 10 {
            from {
                color 128
            }
            action {
                action-type accept
                flex-algo 128
            }
        }
        default-action {
            action-type accept
            flex-algo 128
        }
    }
}
```

Example: classic CLI

```
A:node-2>config>router>policy-options# info
-----
    policy-statement "ExamplePolicy"
        entry 10
            from
                color 128
            exit
            action accept
```

```

        flex-algo 128
    exit
    exit
    default-action accept
        flex-algo 128
    exit
exit

```

After a policy-statement is created it can be applied to the BGP service using service-based import policies. BGP-based automated flexible algorithm steering for SR-MPLS based segment routing can be applied for VPRN, EVP and VPLS services.

This example provides the various places a policy-statement can be applied to initiate BGP-based automated steering to a flexible algorithm:

Example: MD-CLI

```

[ex: /configure service]
A:admin@node-2# info
vpn-apply-import true
  ebgp-default-reject-policy {
    import false
    export false
  }
  import {
    policy ["ExamplePolicy"]
  }
  next-hop-resolution {
    shortcut-tunnel {
      family ipv4 {
        allow-flex-algo-fallback true
      }
    }
    labeled-routes {
      transport-tunnel {
        family vpn {
          allow-flex-algo-fallback true
        }
        family label-ipv4 {
          allow-flex-algo-fallback true
        }
        family label-ipv6 {
          allow-flex-algo-fallback true
        }
      }
    }
  }
}

```

Example: classic CLI

```

A:node-2>config>router>bgp# info
-----
vpn-apply-import
import "ExamplePolicy"
next-hop-resolution
  shortcut-tunnel
    family ipv4
      allow-flex-algo-fallback
      resolution disabled
  exit

```

```

        exit
        labeled-routes
        transport-tunnel
        family vpn
        resolution filter
        allow-flex-algo-fallback
        exit
        family label-ipv4
        resolution filter
        allow-flex-algo-fallback
        exit
        family label-ipv6
        resolution filter
        allow-flex-algo-fallback
        exit
    exit
exit
no shutdown

```

Example: MD-CLI

```

[ex: /configure service]
A:admin@node-2# info
  epipe "3" {
    customer "1"
    bgp 1 {
      vsi-import ["ExamplePolicy"]
    }
  }
  vpls "2" {
    customer "1"
    bgp 1 {
      vsi-import ["ExamplePolicy"]
    }
  }
  vprn "1" {
    customer "1"
    bgp-ipvpn {
      mpls {
        vrf-import {
          policy ["ExamplePolicy"]
        }
        auto-bind-tunnel {
          allow-flex-algo-fallback true
        }
      }
    }
  }
}

```

Example: classic CLI

```

A:node-2>config>service# info
-----
  customer 1 name "1" create
    description "Default customer"
  exit
  vprn 1 name "1" customer 1 create
    shutdown
    bgp-ipvpn
    mpls
      shutdown

```

```

        auto-bind-tunnel
        allow-flex-algo-fallback
    exit
    vrf-import "ExamplePolicy"
exit
exit
exit
vpls 2 name "2" customer 1 create
    shutdown
    bgp
        vsi-import "ExamplePolicy"
    exit
    stp
        shutdown
    exit
exit
epipe 3 name "3" customer 1 create
    shutdown
    bgp
        vsi-import "ExamplePolicy"
    exit
exit

```

One of the side effects of using a flex-algorithm import policy-statement is packets are dropped when no igp-shortcut exists to the BGP next-hop. This behavior ensures that committed Service Level Agreements (SLA) are kept when flexible algorithms are used. However, for some use-cases, dropping packets may be too strict of a behavior, and therefore Nokia SR OS allows through the configuration of the **allow-flex-algo-fallback** command a relaxation of the requirement for a matching flexible algorithm igp-shortcut. When the **allow-flex-algo-fallback** command is configured and no matching shortcut exists, an igp-shortcut from algorithm uses a lower preferred alternative, and therefore, may be breaking strictly committed SLAs.

2.2 Establishing segment routing TE LSPs

When segment routing is used together with MPLS data plane, the SID is a standard MPLS label. A router forwarding a packet using segment routing therefore pushes one or more MPLS labels.

Segment routing using MPLS labels can be used in both shortest path routing applications (see [Segment routing in shortest path forwarding](#) for more information) and in traffic engineering (TE) applications, as described in this section.

An SR-TE LSP supports a primary path, with FRR backup, and one or more secondary paths. A secondary path can be configured as standby.

SR OS implements the following computation methods for the paths of an SR-TE LSP:

- **hop-to-label translation**

The TE-DB converts the list of hops, destination of the LSP, and the strict or loose hops in the path definition to a list of SIDs by searching the IGP instances with segment routing enabled. This method does not support TE constraints except for loose or strict hops.

See [SR-TE LSP path computation using hop-to-label translation](#) for more information.

- **local CSPF**

The LSP path TE constraints are considered in the path computation. This method implements most of the CSPF capabilities supported with RSVP-TE LSP with very few exceptions, such as the bandwidth constraint which cannot be reserved with SR-TE LSP because of the lack of a signaling protocol to establish the LSP path.

See [SR-TE LSP path computation using local CSPF](#) for more details.

- **Path Computation Element (PCE)**

The router acting as a PCE client (PCC) requests a computation of the path of an SR-TE LSP from the PCE using the PCEP.



Note: PCEP is not supported on the 7705 SAR Gen 2.

- **user-specified SID list**

The SR-LSP feature provides the option for the user to manually configure each path of the LSP using an explicit list of SID values.

See [SR-TE LSP paths using explicit SIDs](#) for more details.

The configured or computed path of an SR-TE LSP can use a combination of node SIDs and adjacency SIDs.

2.2.1 SR-TE MPLS support

The following MPLS commands and nodes are supported in Segment Routing-Traffic Engineering (SR-TE):

- global MPLS-level commands and nodes

interface, lsp, path, shutdown

- LSP-level commands and nodes

bfd, bgp-shortcut, bgp-transport-tunnel, cspf, exclude, hop-limit, igp-shortcut, include, metric, metric-type, path-computation-method, primary, retry-limit, retry-timer, revert-timer, shutdown, to, from, vprn-auto-bind

- Both primary and secondary paths are supported with an SR-TE LSP. The following primary path level commands and nodes are supported with SR-TE LSP:

bandwidth, bfd, delay-metric-limit, exclude, hop-limit, include, priority, shutdown

The following secondary path level commands and nodes are supported with SR-TE LSP:

bandwidth, bfd, delay-metric-limit, exclude, hop-limit, include, path-preference, priority, shutdown, srlg, standby

The following MPLS commands and nodes are not supported with SR-TE LSP:

- global MPLS-level commands and nodes (configuration ignored)

admin-group-frr, auto-bandwidth-multipliers, auto-lsp, bypass-resignal-timer, cspf-on-loose-hop, dynamic-bypass, exponential-backoff-retry, frr-object, hold-timer, ingress-statistics, least-fill-min-thd, least-fill-reoptim-thd, logger-event-bundling, lsp-init-retry-timeout, lsp-template, max-bypass-associations, mbb-prefer-current-hops, mpls-tp, p2mp-resignal-timer, p2mp-s2l-fast-retry, p2p-active-path-fast-retry, retry-on-igp-overload, secondary-fast-retry-timer, shortcut-local-ttl-propagate, shortcut-transit-ttl-propagate, srlg-database, srlg-frr, static-lsp, static-lsp-fast-retry, user-srlg-db

- LSP-level commands and nodes not supported (configuration blocked)

adaptive, adspec, auto-bandwidth, class-type, dest-global-id, dest-tunnel-number, exclude-node, fast-reroute, ldp-over-rsvp, least-fill, main-ct-retry-limit, p2mp-id, primary-p2mp-instance, propagate-admin-group, protect-tp-path, rsvp-resv-style, working-tp-path

- primary path level commands and nodes not supported (configuration blocked)

adaptive, backup-class-type, class-type, record, record-label

- secondary path level commands and nodes not supported (configuration blocked)

adaptive, class-type, record, record-label

The user can associate an empty path or a path with strict or loose explicit hops with SR-TE LSP paths using the **hop**, **primary**, and **secondary** CLI commands.

A hop that corresponds to an adjacency SID must be identified with its far-end host IP address (next hop) on the subnet. If the local end host IP address is provided, this hop is ignored because this router can have multiple adjacencies (next hops) on the same subnet.

A hop that corresponds to a node SID is identified by the prefix address.

2.2.2 SR-TE LSP path computation

The path is computed using the hop-to-label translation method. MPLS makes a request to the TE-DB to get the label corresponding to each hop entered by the user in the primary path of the SR-TE LSP.

The user can configure the maximum number of labels that the ingress LER can push for a specified SR-TE LSP by using the **max-sr-labels** command. This command sets a limit on the maximum label stack size of the SR-TE LSP primary path, which allows room to insert additional transport, service, and other labels when packets are forwarded in a specific context.

Use the **config>router>mpls>lsp>max-sr-labels label-stack-size [additional-frr-labels labels]** CLI command to configure the maximum number of labels.

Set the **max-sr-labels label-stack-size** value to account for the required maximum label stack of the primary path of the SR-TE LSP.

The 7705 SAR Gen 2 in the ILER role supports pushing a maximum of 11 four-byte labels or outermost ETH-VLAN tags. This maximum label stack includes all of the following label types:

- service
- CW
- hash
- BGP-LU
- transport
- LFA
- TI-LFA
- network egress
- dot1q VLAN

Set the **additional-frr-labels labels** value to account for additional labels inserted by remote LFA or Topology Independent LFA (TI-LFA) for the backup next-hop of the SR-TE LSP. The supported range is 0 to 3 labels with a default value of 1.

The sum of the value of both labels represents the worst-case transport of the SR label stack size for this SR-TE LSP and is populated by MPLS in the TTM. Services can check the value to decide if a service can be bound or a route can be resolved to this SR-TE LSP. See [SR-TE label stack check for services and shortcuts](#) for more information about the label stack size check and requirements for services and shortcut applications.

2.2.3 SR-TE LSP path computation using hop-to-label translation

MPLS passes the path information to the TE-DB, which converts the list of hops into a label stack as follows:

- A loose hop with an address matching any interface (loopback or not) of a router (identified by the router ID) is always translated to a node SID. If the prefix matching the hop address has a node SID in the TE database, it is selected by preference. If not, the node SID of any loopback interface of the same router that owns the hop address is selected. In the latter case, the lowest IPv4(IPv6) address of that router that has a /32 (/128) prefix SID is selected.
- A strict hop with an address matching any interface (loopback or not) of a router (identified by the router ID) is always translated to an adjacency SID. If the hop address matches the host address reachable in a local subnet from the previous hop, then the adjacency SID of that adjacency is selected. If the hop address matches a loopback interface, it is translated to the adjacency SID of any link from the previous hop which terminates on the router owning the loopback. The adjacency SID label of the selected link is used.

In both cases, it is possible to have multiple matching previous hops in the case of a LAN interface. In this case, the adjacency-SID with the lowest interface address is selected.

- In addition to the IGP instance that resolved the prefix of the destination address of the LSP in the RTM, all IGP instances are scanned from the lowest to the highest instance ID, beginning with IS-IS instances and then OSPF instances. For the first instance via which all specified path hop addresses can be translated, the label stack is selected. The hop-to-label translation tool does not support paths that cross area boundaries. All SIDs and labels of a path are therefore taken from the same IGP area and instance.
- Unnumbered network IP interfaces, which are supported in the router's TE database, can be selected when converting the hops into an adjacency SID label when the user has entered the address of a loopback interface as a strict hop; however, the user cannot configure an unnumbered interface as a hop in the path definition.



Note: For the hop-to-label translation to operate, the user must enable TE on the network links, which means adding the network interfaces to MPLS and RSVP. The user must also enable the **traffic-engineering** option on all participating router IGP instances. If any router has the **database-export** option enabled in the participating IGP instances to populate the learned IGP link state information into the TE-DB, then enabling the **traffic-engineering** option is not required. For consistency purposes, Nokia recommends that the **traffic-engineering** option is always enabled.

2.2.4 SR-TE LSP path computation using local CSPF

This section describes full CSPF path computation for SR-TE LSP paths.

The **path-computation-method [local-cspf]** command configures the path computation method for SR-TE LSPs. It enables the user to select the computation method for the SR-TE LSP and set it to hop-to-label translation or local CSPF path computation method. The **no** form of this command, which is the default value, sets the computation method to the hop-to-label translation method.

2.2.4.1 Extending MPLS and TE database CSPF support to SR-TE LSP

The following MPLS and TE database features extend CSPF support to SR-TE LSP:

- IPv4 SR-TE LSP
- local CSPF on both primary and secondary standby paths of an IPv4 SR-TE LSP
- local CSPF in LSP templates of types **mesh-p2p-srte** and **one-hop-p2p-srte** of SR-TE auto-LSP
- support path computation in single area OSPFv2 and IS-IS IGP instances
- computes full explicit TE paths using TE links as hops and returning a list of SIDs consisting of adjacency SIDs and parallel adjacency set SIDs. SIDs of a non-parallel adjacency set are not used in CSPF. The details of the CSPF path computation are provided in [SR-TE specific TE-DB changes](#). Loose-hop paths, using a combination of node SID and adjacency SID, are not required.
- use random path selection in the presence of ECMP paths that satisfy the LSP and path constraints. Least-fill path selection is not required.
- provide an option to reduce or compress the label stack such that the adjacency SIDs corresponding to a segment of the explicit path are replaced with a node SID whenever the constraints of the path are met by all the ECMP paths to that node SID. The details of the label reduction are provided in [SR-TE LSP path label stack reduction](#).
- use legacy TE link attributes as in RSVP-TE LSP CSPF
- uses timer reoptimization of all paths of the SR-TE LSP that are in the operational up state. This differs from the RSVP-TE LSP resignal timer feature which re-optimizes the active path of the LSP only.

MPLS provides the current path of the SR-TE LSP and TE-DB updates the total IGP or TE metric of the path, checking the validity of the hops and labels as per current TE-DB link information. CSPF then calculates a new path and provides both the new and metric updated current path back to MPLS. MPLS programs the new path only if the total metric of the new computed path is different from the updated metric of the current path, or if one or more hops or labels of the current path are invalid. Otherwise, the current path is considered one of the most optimal ECMP paths and is not updated in the datapath.

Timer resignal applies only to the CSPF computation method and not to the IP-to-label computation method.

- use manual reoptimization of a path of the SR-TE LSP. In this case, the new computed path is always programmed even if the metric or SID list is the same.
- supports ad hoc reoptimization. This feature triggers the ad hoc resignaling of all SR-TE LSPs if one or more IGP link down events are received in TE-DB. For more information, see [Ad hoc SR-TE LSP reoptimization on receipt of IGP link events](#).
- support unnumbered interfaces in the path computation. There is no support for configuring an unnumbered interface as a hop in the path of the LSP. The path can be empty or include hops with the address of a system or loopback interface, but path computation can return a path that uses TE links corresponding to unnumbered interfaces.
- support **admin-group**, **hop-count**, IGP metric, and TE-metric constraints

- bandwidth constraint is not supported because SR-TE LSP does not have an LSR state to book bandwidth. The **bandwidth** parameter, when enabled on the LSP path, has no impact on local CSPF path calculation. However, the **bandwidth** option is passed to PCE when it is the selected path computation method. PCE reserves bandwidth for the SR-TE LSP path accordingly.

2.2.4.2 SR-TE specific TE-DB changes

When the **traffic-engineering** command is enabled in an OSPFv2 instance, only local and remote TE-enabled links are added into the TE-DB. A TE-link is a link that has one or more TE attributes added to it in the MPLS interface context. Link TE attributes are TE metric, bandwidth, and membership in an SRLG or an Admin-Group.

To allow the SR-TE LSP path computation to use SR-enabled links that do not have TE attributes, the following changes are made:

- OSPFv2 passes all links, whether they are TE-enabled or SR-enabled, to the TE-DB, as currently performed by IS-IS.
- TE-DB relaxes the link back-check when performing a CSPF calculation to ensure that there is at least one link from the remote router to the local router. Because OSPFv2 advertises the remote link IP address or remote link identifier only when a link is TE-enabled, the strict check about the reverse direction of a TE-link cannot be performed if the link is SR-enabled but not TE-enabled.

As a consequence of this implementation, CSPF can compute an SR-TE LSP with SR-enabled links that do not have TE attributes. This means that if the user admin shuts down an interface in MPLS, an SR-TE LSP path that uses this interface does not go operationally down.

2.2.4.3 SR-TE LSP and auto-LSP-specific CSPF changes

The local CSPF for an SR-TE LSP is performed in two phases. Phase 1 computes a fully explicit path with all TE links to the destination specified, as in the case of an RSVP-TE LSP. If the user has enabled label stack reduction or compression for this LSP, Phase 2 is applied to reduce the label stack so that adjacency SIDs corresponding to a segment of the explicit path are replaced with a node SID whenever the constraints of the path are met by all the ECMP paths to that node SID. The details of the label reduction are provided in [SR-TE LSP path label stack reduction](#).

The CSPF computation algorithm for the fully explicit path in Phase 1 remains mostly unchanged from its behavior in RSVP-TE LSP.

The meaning of a strict and loose hop in the path of the LSP is the same as in CSPF for RSVP-TE LSP. A strict hop means that the path from the previous hop must be a direct link. A loose hop means the path from the previous hop can traverse intermediate routers.

A loose hop may be represented by a set of back-to-back adjacency SIDs if not all paths to the node SID of that loose hop satisfy the path TE constraints. This is different from the IP-to-labeled path computation method where a loose hop always matches a node SID because no TE constraints are checked in the path to that loose hop.

When the label stack of the path is reduced or compressed, a strict hop may be represented by a node SID if all the links from the previous hop satisfy the path TE constraints. This is different from the IP-to-labeled path computation method where a strict hop always matches an adjacency SID or a parallel adjacency set SID.

The first phase of CSPF returns a full explicit path with each TE link specified all the way to the destination. The label stack may contain protected adjacency SIDs, unprotected adjacency SIDs, and adjacency set

SIDs. The user can configure the type of adjacency protection for the SR-TE LSP using a CLI command as described in [SR-TE LSP path protection](#).

SR OS does not support the origination of a global adjacency SID. If received from a third-party router implementation, it is added into the TE database but is not used in any CSPF path computation.

2.2.4.3.1 SR-TE LSP path protection

SR-TE LSP allows the user to configure whether the path of the LSP must use protected or unprotected adjacencies exclusively for all links of the path.

When SR OS routers form an IGP adjacency over a link and segment-routing context is enabled in the IGP instance, the static or dynamic label assigned to the adjacency is advertised in the link adjacency SID sub-TLV. By default, an adjacency is always eligible for LFA/RLFA/TI-LFA protection and the B-flag in the sub-TLV is set. The presence of a B-flag does not reflect the instant state of the availability of the adjacency LFA backup; it reflects that the adjacency is eligible for protection. The SR-TE LSP using the adjacency in its path still comes up if the adjacency does not have a backup programmed in the datapath at that instant. Use the **configure>router>isis>interface> no sid-protection** command to disable protection. When protection is disabled, the B-flag is cleared and the adjacency is not eligible for protection by LFA/RLFA/TI-LFA.

SR OS also supports the adjacency set feature that treats a set of adjacencies as a single object and advertises a link adjacency sub-TLV for it with the S-flag (SET flag) set to 1. The adjacency set in the SR OS implementation is always unprotected, even if there is a single member link in it and therefore the B-flag is always clear. Only a parallel adjacency set, meaning that all links terminate on the same downstream router, is used by the local CSPF feature.

The same P2P link can participate in a single adjacency and in one or more adjacency sets. Therefore, multiple SIDs can be advertised for the same link.

Third party implementations of Segment Routing may advertise two SIDs for the same adjacency when LFA is enabled in the IS-IS or OSPF instance: one protected with the B-flag set and one unprotected with the B-flag clear. SR OS can achieve the same behavior using one of the following two options:

- Enabling the allocation of dual SIDs using the following command for IS-IS or OSPF respectively:

- **MD-CLI**

```
configure router isis segment-routing adjacency-sid allocate-dual-sids true
configure router ospf segment-routing adjacency-sid allocate-dual-sids true
```

- **classic CLI**

```
configure router isis segment-routing adjacency-sid allocate-dual-sids
configure router ospf segment-routing adjacency-sid allocate-dual-sids
```

- Adding a link to a single-member adjacency SET, in which case a separate SID is advertised for the SET and the B-flag is cleared while the SID for the regular adjacency over that link has its B-flag set by default



Note: LFA must be enabled under the IS-IS or OSPF instance for the previously mentioned cases.

In all cases, SR OS CSPF can use all local and remote SIDs to compute a path for an SR-TE LSP based on the needed local protection property.

The following different behaviors of CSPF are introduced with SR-TE LSP:

- If the **local-sr-protection** command is not enabled (**no local-sr-protection**) or is set to **preferred**, the local CSPF prefers a protected adjacency over an unprotected adjacency whenever both exist for a TE link. This is done on a link-by-link basis after the path is computed based on the LSP path constraints. This means that the protection state of the adjacency is not used as a constraint in the path computation. It is only used to select an SID among multiple SIDs after the path is selected. Thus, the computed path can combine both types of adjacencies.

If a parallel adjacency set exists between two routers in a path and all the member links satisfy the constraints of the path, it is selected in preference to a single protected adjacency, which is selected in preference to a single unprotected adjacency.

If multiples ECMP paths satisfy the constraints of the LSP path, one path is selected randomly and then the SID selection above applies. There is no check if the selected path has the highest number of protected adjacencies.

- If the **local-sr-protection** command is set to a value of **mandatory**, CSPF uses it as an additional path constraint and selects protected adjacencies exclusively in computing the path of the SR-TE LSP. Adjacency sets cannot be used because they are always unprotected.

If no path satisfies the other LSP path constraints and consists of all TE links with protected adjacencies, the path computation returns no path.

- If the **local-sr-protection** command is set to a value of **none**, CSPF uses it as an additional path constraint and selects unprotected adjacencies exclusively in computing the path of the SR-TE LSP.

If a parallel adjacency set exists between two routers in a path and all the member links satisfy the constraints of the path, it is selected in preference to a single unprotected adjacency.

If no path satisfies the other LSP path constraints and consists of all TE links with unprotected adjacencies, the path computation returns no path.

The **local-sr-protection** command impacts PCE-computed and PCE-controlled SR-TE LSPs. When the **local-sr-protection** command is set to the default value **preferred**, or to the explicit value of **mandatory**, the **local-protection-desired** flag (L-flag) in the LSPA object in the PCReq (Request) message or in the PCRpt (Report) message is set to a value of 1.

When the **local-sr-protection** command is set to **none**, the **local-protection-desired** flag (L-flag) in the LSPA object is cleared. The PCE path computation checks this flag to decide if protected adjacencies are used in preference to unprotected adjacencies (L-flag set) or must not be used at all (L-flag clear) in the computation of the SR-TE LSP path.

2.2.4.3.2 SR-TE LSP path label stack reduction

The objective of the label stack reduction is twofold:

- It reduces the label stack so ingress PE routers with a lower Maximum SID Depth (MSD) can still work.
- It provides the ability to spray packets over ECMP paths to an intermediate node SID when all these paths satisfy the constraints of the SR-TE LSP path. Even if the resulting label stack is not reduced, this aspect of the feature is still useful.

If the user enables the **label-stack-reduction** command for this LSP, a second phase is applied, attempting to reduce the label stack that resulted from the fully explicit path with adjacency SIDs and adjacency sets SIDs computed in the first phase.

This is to attempt a replacement of adjacency and adjacency set SIDs corresponding to a segment of the explicit path with a node SID whenever the constraints of the path are met by all the ECMP paths to that node SID.

The label stack reduction algorithm uses the following procedure.

1. Phase 1 of the CSPF returns up to three fully explicit ECMP paths that are eligible for label stack reduction. These paths are equal cost from the point of view of IGP metric or TE metric as configured for that SR-TE LSP.
2. Each fully explicit path of the SR-TE LSP that is computed in Phase 1 of the CSPF is split into a number of segments that are delimited by the user-configured loose or strict hops in the path of the LSP. Label stack reduction is applied to each segment separately.
3. Label stack reduction in Phase 2 consists of traversing the CSPF tree for each ECMP path returned in Phase 1, then attempting to find the farthest node SID in a path segment that can be used to summarize the entire path up to that node SID. This requires that all links of ECMP paths are able to reach the node SID from the current node on the CSPF tree to satisfy all the TE constraints of the SR-TE LSP paths. ECMP is based on the IGP metric, in this case, because this is what routers use in the datapath when forwarding a packet to the node SID.

If the TE metric is enabled for the SR-TE LSP, one of the constraints is that the TE metric must be the same value for all the IGP metric ECMP paths to the node SID.

4. CSPF in Phase 2 selects the first candidate ECMP path from Phase 1, which reduced label stack that satisfies the constraint carried in the **max-sr-labels** command.
5. The CSPF path computation in Phase 1 always avoids a loop over the same hop, as is the case with the RSVP-TE LSP. In addition, the label stack reduction algorithm prevents a path from looping over the same hop because of the normal routing process. For example, it checks if the same node is involved in the ECMP paths of more than one segment of the LSP path and builds the label stack to avoid this situation.
6. During the MBB procedure of a timer or the manual re-optimization of an SR-TE LSP path, the TE-DB performs the following steps in addition to the initial path computation.
 - MPLS provides the TE-DB with the current working path of the SR-TE LSP.
 - The TE-DB updates the path's metric based on the IGP or TE link metric (if the TE metric enabled for the SR-TE LSP).
 - For each adjacency SID, the TE-DB verifies that the related link and SID are still in its database and that the link fulfills the LSP constraints. If so, it picks up the current metric.
 - For each node SID, the TE-DB verifies that the related prefix and SID are still available, and if so, checks that all the links on the shortest IGP path to the node owning the node SID fulfill the SR-TE LSP path constraints. This step reuses the same checks detailed in step 3 for the label compression algorithm.
 - CSPF computes a new path with or without label stack reduction as described in steps 1, 2, and 3.
 - The TE-DB returns both paths to MPLS. MPLS always programs the new path in the case of a manual re-optimization. MPLS compares the metric of the new path to the current path and if different, programs the new path in the case of a timer re-optimization.
7. The TE-DB sends the reduced path ERO and label stack to MPLS, along with the following information:
 - a list of SRLGs of each hop in the ERO, represented by a node SID, including the SRLGs used by links in all ECMP paths to reach that node SID from the previous hop

- the cost of each hop in the ERO represented by an adjacency SID or adjacency set SID. The cost corresponds to the IGP metric or TE metric (if the TE metric is enabled for the SR-TE LSP) of that link or set of links. In the case of an adjacency set, all TE metrics of the links must be the same, otherwise CSPF does not select the set.
 - the cost of each hop in the ERO represented by a node SID, which corresponds to the cumulated IGP metric or TE metric (if the TE metric is enabled for the SR-TE LSP) to reach the node SID from the previous hop using the fully explicit path computed in Phase 1.
 - the total cost or computed metric of the SR-TE LSP path. This consists of the cumulated IGP metric or TE metric (if TE metric enabled for the SR-TE LSP) of all hops of the fully explicit path computed in Phase 1 of the CSPF.
8. If label stack reduction is disabled, the values of the **max-sr-labels** and the **hop-limit** commands are applied to the full explicit path in Phase 1.
- The minimum of the two values is used as a constraint in the full explicit path computation.
- If the resulting ECMP paths net hop-count in Phase 1 exceeds this minimum value, the TE-DB does not return a path to MPLS.
9. If label stack reduction is enabled, the values of the **max-sr-labels** and the **hop-limit** commands are both ignored in Phase 1 and only the value of the **max-sr-labels** command is used as a constraint in Phase 2.
- If the Phase 2 reduction of all candidate paths results in a net label stack size that exceeds the value of the **max-sr-labels** command, the TE-DB does not return a path to MPLS.
10. The label stack reduction uses a node SID to replace a segment of the SR-TE LSP path; using an anycast SID or a prefix SID with the N-flag clear is not supported.

2.2.4.3.3 Interaction with SR-TE LSP path protection

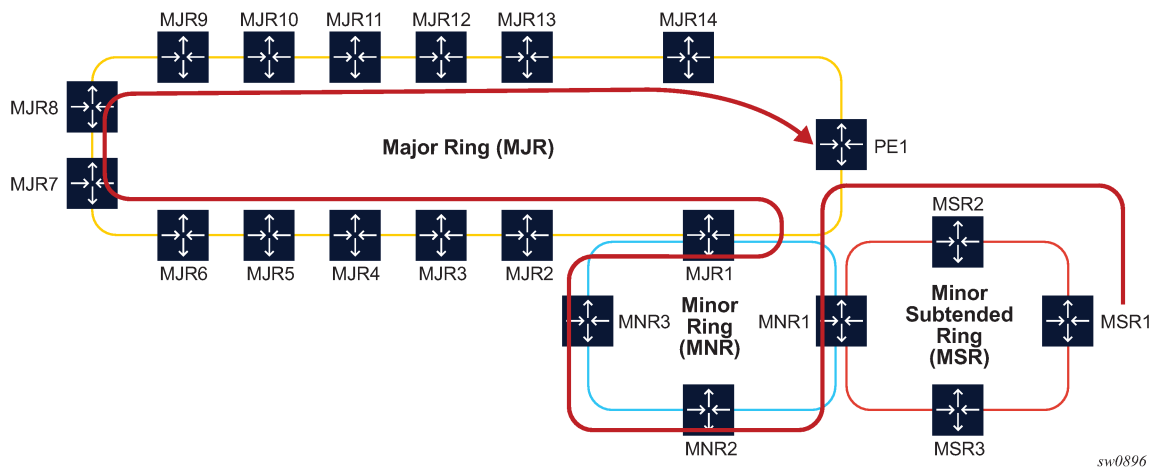
Label stack reduction is only attempted when the path protection **local-sr-protection** command is disabled or configured to the value of **preferred**.

If **local-sr-protection** is configured to a value of **none** or **mandatory**, the command is ignored, and the fully explicit path computed in Phase 1 is returned by the TE-DB CSPF routine to MPLS. This is because a node SID used to replace an adjacency SID or an adjacency set SID can be unprotected or protected by LFA and this is based on local configuration on each router which resolves this node SID but is not directly known in the information advertised into the TE-DB. Therefore, CSPF cannot enforce the protection constraint requested along the path to that node SID.

2.2.4.3.4 Examples of SR-TE LSP path label stack reduction

The following figure shows a metro aggregation network with three levels of rings for aggregating regional traffic from edge ring routers to a PE router.

Figure 19: Label stack reduction in a 3-tier ring topology



The path of the highlighted LSP uses admin groups to force the traffic eastwards or westwards over the 3-ring topologies such that it uses the longest path possible. Assume all links in a bottom-most ring1 have admin-group=east1 for the eastward direction and admin-group=west1 for the westward direction.

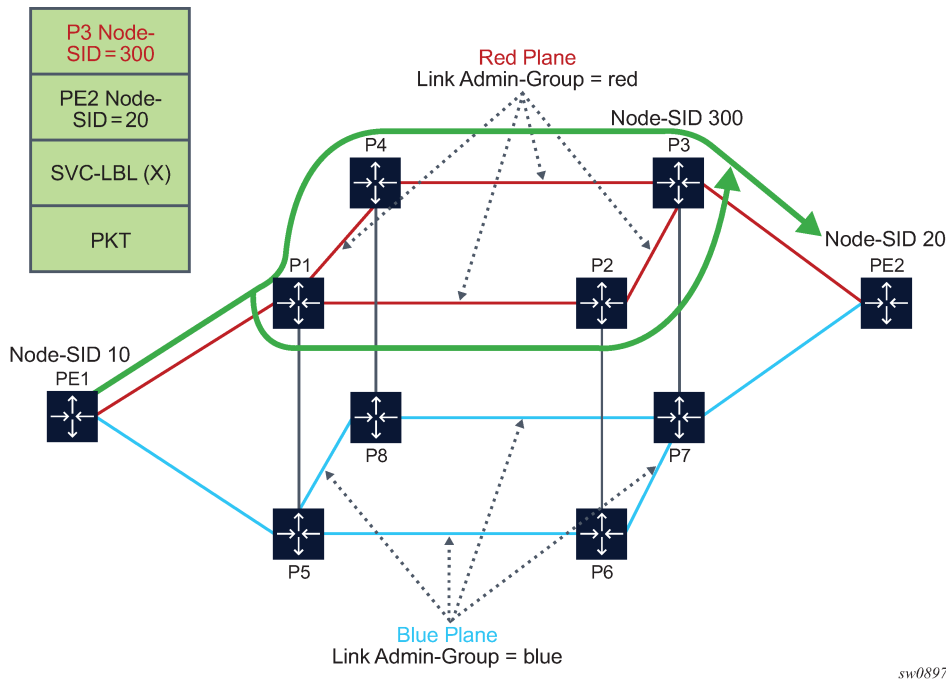
Similarly, links in middle ring2 have admin-group=east2 and admin-group=west2, and links in top-most ring3 have admin-group=east3 and admin-group=west3. To achieve the longest path shown, the LSP or path should have an include statement: include east1 west2 east3. The fully explicit path computed in Phase 1 of CSPF results in label stack of size 18.

The label stack reduction algorithm searches for the farthest node SID in that path, which can replace a segment of the strict path while maintaining the stated admin-group constraints. The reduced label stack contains the SID adjacency MSR1-MSR2, the found node SIDs plus the node SID of the destination for a total of four labels to be pushed on the packet (the label for the adjacency MSR1-MSR2 is not pushed):

{N-SID MNR2, N-SID of MNR3, N-SID of MJR8, N-SID of PE1}

The following figure shows an example topology that creates two TE planes by applying a common admin group to all links of a specified plane. There are a total of four ECMP paths to reach PE2 from PE1, two within the red plane and two within the blue plane.

Figure 20: Label stack reduction in the presence of ECMP paths



For an SR-TE LSP from PE1 to PE2, which includes the red admin-group as a constraint, Phase 1 of CSPF results in two fully explicit paths using adjacency SID of the red TE links:

path 1 = {PE1-P1, P1-P2, P2-P3, P3-PE2}

path 2 = {PE1-P1, P1-P4, P4-P3, P3-PE2}

Phase 2 of CSPF finds node SID of P3 as the farthest hop it can reach directly from PE1 while still satisfying the constraint to include the red admin-group constraint. If the node SID of PE2 is used as the only SID, then traffic would also be sent over the blue links.

Then, the reduced label stack is: {P3 Node-SID=300, PE2 Node-SID=20}.

The resulting SR-TE LSP path combines the two explicit paths out of Phase 1 into a single path with ECMP support.

2.2.4.4 Delay metric

SR OS supports the **delay** metric type option for LSPs and LSP templates.



Note: To advertise network delay information, use the following commands:

- **configure router isis traffic-engineering-options advertise-delay**
- **configure router isis traffic-engineering-options application-link-attributes**
- **configure router isis traffic-engineering-options application-link-attributes legacy**

See [Support of MT-ISIS MT2 TE link attributes and router capability](#) for more information.

When the **delay** metric type is configured, SR OS selects the path throughout the network with the lowest latency that is below the configured delay metric limit under the primary path or secondary path of the LSP.

If the delay metric limit is not configured and the **delay** metric type is configured, the system computation selects the path with the lowest end-to-end delay. If the delay metric limit is configured and if the latency does not satisfy that limit, the path is not set up.

SR OS can set the delay metric within a network in one of the following ways:

- statically, under the interface itself
- dynamically, using TWAMP light

Each path within the LSP can be configured with its own delay metric limit. This is useful when different redundant paths within the LSP are going through different networks (for example, different operators) with different latency values. Each path can be configured with its own delay metric limit associated with its corresponding network.

2.2.4.4.1 Delay metric implementation notes

The following are implementation notes for the delay metric:

- The delay measurement is accumulative and end to end.
- The metric type can be configured as **delay** only if:
 - the path computation method is configured as **local-cspf**
 - PCE control is disabled
- The system computation uses the delay metric limit only when the metric type is configured as **delay**. Otherwise, the configured limit is not used in computation.
- If PCE reporting is enabled, the LSP with the metric type configured as **delay** is reported to the PCE controller as metric type **te**. See "TE metric (IS-IS and OSPF)" in the *7705 SAR Gen 2 MPLS Guide* MPLS Guide for more information about the **te** option.
- When the metric type is configured as **delay**, the following applies:
 - The LSP path operational metric (for example, the LSP metric to be used for IGP shortcuts) is set to the maximum metric value, that is, 16777215.
 - Label stack reduction does not function. Label stack reduction can result in the usage of links with a high delay. Nokia recommends avoiding the combination of label stack reduction and metric type **delay**.
- The timer-based reoptimization functionality is the same for metric type **delay**, **igp**, or **te**.
- Changes in link delay are not supported as an IGP event triggering the ad-hoc reoptimization of SR-TE LSPs.
- When modifying the delay metric limit configurations on the LSP path, the system brings down the LSP path, then retries the path. If there is no CSPF path that meets the limit, the LSP does not come back up.
- When the LSP is down, the LSP tries the delay constraint and tries to establish when the retry timer expires.
- When the LSP is up and the delay becomes higher than the configured limit or if there is a change in the IGP delay configuration, the LSP does not go down unless it is explicitly cleared or bounced. Even if resignaling occurs, the LSP remains up; however, the MBB state reflects a failure.
- If the primary path delay increases beyond the delay metric limit and the secondary path delay is below the limit, the primary path does not switch to the secondary path. The primary path stays active until it

is operationally down, then switches to the secondary path. There is no trap or event indication on the primary path that the path delay has become higher than the limit.

- For IS-IS, the **delay** metric type is supported for multitopology (MT) 0 and 2.

2.2.4.4.2 Delay metric support

The delay metric is supported on the following:

- segment routing LSPs and SR-TE LSPs
- static LSP configurations where the **delay** option can be configured for each LSP and each path of the LSP can be configured with its own delay metric limit. If the LSP metric type is **delay** and the limit is not configured, the system computation selects the lowest latency path.
- LSP templates. The delay metric limit applies to the entire template.

2.2.4.5 Ad hoc SR-TE LSP reoptimization on receipt of IGP link events

The following command enables the ad hoc reoptimization of all CSPF paths in operational up state of all SR-TE LSPs at the receipt of an IGP link event.

```
configure router mpls sr-te-resignal resignal-on-igp-event
```

The following link events are supported:

- link down
- link up
- IGP or TE metric change
- SRLG change
- admin group change

The ad hoc reoptimization follows the same behavior as in the timer based resignal Make-Before-Break (MBB) feature. MPLS re-evaluates all the paths in the operational up state of all SR-TE LSPs. The re-evaluation consists of updating the total IGP or TE metric of the current path, checking the validity of the hops and labels, and computing a new CSPF path. MPLS programs the new path only if its total metric is different from the updated metric of the current path, or if one or more hops or labels of the current path are invalid. Otherwise, the current path is considered to be the most optimal and retained.

This feature does not require that the timer-based resignal command, as follows, be enabled. If enabled, it ends the resignal timer and performs the ad hoc reoptimization.

```
configure router mpls sr-te-resignal resignal-timer
```

2.2.5 SR-TE LSP paths using explicit SIDs

SR OS supports the ability for SR-TE primary and secondary paths to use a configured path containing explicit SID values. The SID value for an SR-TE LSP hop is configured using the **sid-label** *sid-value* parameter of the **configure>router mpls path hop** command, where *sid-value* specifies an MPLS label value for that hop in the path.

When SIDs are explicitly configured for a path that consists of either all SIDs or all IP address hops, the user must provide all of the necessary SIDs to reach the destination. The router does not validate whether the provided label stack is correct.

A path containing SID label hops is used even if **path-computation-method {local-cspf | pce}** is configured for the LSP. That is, the path computation method configured at the LSP level is ignored when explicit SIDs are used in the path. This means that the router can bring up the path if the configured path contains SID hops even if the LSP has path computation enabled.



Note: When an LSP consists of some SID labeled paths and some paths under local-CSPF computation, the router cannot guarantee SRLG diversity between the CSPF paths and the SID labeled paths. CSPF is not aware of the existence of the SID labeled paths because they are not listed in the TE database.

Paths containing explicit SID values can only be used by SR-TE LSPs.

2.2.6 SR-MPLS IGP shortcuts over SR-TE LSP

SR OS supports resolving an SR-MPLS shortest path tunnel (SR-ISIS and SR-OSPF) over IGP shortcuts using SR-TE LSPs. By default, the SR-TE LSPs are not eligible as SR-MPLS IGP shortcuts. This configuration reduces the risk of accidental SR-MPLS forwarding loops.

To enable SR-MPLS IGP shortcuts over SR-TE LSP, configure SR-TE LSPs to be eligible as IGP shortcuts and configure IGP so that SR-TE LSPs can be used as IGP shortcuts in SR-MPLS.

Use the following commands to make an SR-TE LSP eligible as a SR-MPLS IGP shortcut:

- **MD-CLI**

```
configure router mpls lsp igp-shortcut allow-sr-over-srte
configure router mpls lsp igp-shortcut relative-metric
```

- **classic CLI**

```
configure router mpls lsp sr-te igp-shortcut [lfa-protect | lfa-only] allow-sr-over-srte
configure router mpls lsp sr-te igp-shortcut relative-metric [offset] allow-sr-over-srte
```

Use the following commands to make IGP consider the eligible SR-TE LSPs as IGP shortcuts in SR-MPLS.

```
configure router ospf igp-shortcut allow-sr-over-srte
configure router isis igp-shortcut allow-sr-over-srte
```

The following considerations apply when utilizing SR-TE LSPs as SR-MPLS IGP shortcuts:

- **SR-TE LSP composition:** A configured SR-TE LSP intended as a shortcut for an SR-MPLS tunnel must consist of a SR-TE path with an explicit segment list comprising adjacency SID labels that describe the complete end-to-end path.
- **SR-TE LSP initiation:** The SR-TE LSP path may be initiated by either a PCC or a PCE. Its segment list may be manually configured, computed by the router local CSPF, or computed by the PCE, provided that its segment list consists exclusively of adjacency SIDs.
- **Usage criteria for SR-TE LSP as IGP shortcut:** An SR-TE LSP is utilized by the IGP as a shortcut for an SR-MPLS tunnel only if the top SID in the SR-TE path list is an adjacency SID or adjacency set. If the top SID is a node SID, the SR-TE LSP is not used as an IGP shortcut for SR-MPLS tunnels but

may still serve as an IGP shortcut for IP routes and LDP (Label Distribution Protocol) FECs (Forwarding Equivalence Classes).

- sBFD deployment recommendation: Nokia recommends deploying sBFD over an SR-TE LSP utilized as an IGP shortcut for an SR-MPLS tunnel.
- Mixing **allow-sr-over-srte** configurations:
 - IP prefixes resolved via IGP shortcuts may include a mixture of SR-TE LSPs with the **allow-sr-over-srte** configuration enabled or disabled as well as IP next hops. However, an SR-TE LSP with **allow-sr-over-srte** explicitly configured is prioritized over all other next hops and is selected as the first available next hop in the list of candidate next hops.
 - SR-MPLS tunnels resolved via IGP shortcuts can only utilize a mixture of SR-TE LSPs that have **allow-sr-over-srte** configured and IP next hops. Consequently, the SR-MPLS next hops may represent a subset of the available next hops of the corresponding IP prefix.
- TTM preference for next hop selection: Similar to IP prefixes resolved via IGP shortcuts, the configured TTM preference determines the prioritization between RSVP shortcuts and SR-TE shortcuts when resolving SR-MPLS next hops.
- ECMP Path Composition Constraints: When SR-TE shortcuts are enabled and multiple SR-TE LSPs are available to resolve a next hop of an SR-MPLS tunnel, the SR-TE shortcut LSP with the least number of labels is preferred. As with IP prefixes resolved via IGP shortcuts, a set of ECMP next hops will not comprise a mixture of RSVP shortcuts and SR-TE shortcuts when resolving SR-MPLS next hops. A set of ECMP next hop can however mix RSVP-TE and IP next hops or SR-TE and IP next hops.
- Top SID Resolution: An SR-TE LSP path or SR policy candidate path with a top SID that further resolves to an IGP shortcut using a SR-TE LSP will result in packet drops.

To verify the eligibility of an SR-TE LSP as an IGP shortcut for SR-MPLS tunnels, use the appropriate command to review the TTM SR-TE tunnel flags.

```
show router tunnel-table protocol sr-te detail
```

Output example: Displaying SR-TE tunnel flags to verify SR-TE LSP eligibility as an SR-MPLS IGP shortcut

```
=====
Tunnel Table (Router: Base)
=====
Destination : 10.20.1.6/32
NextHop : 1.0.13.1 (524291, ospf (0))
Tunnel Flags : is-over-tunnel entropy-label-capable allow-sr-over-sr-te
Age : 00h01m12s
CBF Classes : (Not Specified)
Owner : sr-te                               Encap : MPLS Tunnel
ID : 655366                                Preference : 8
Tunnel Label : 524287                       Tunnel Metric : 200
Tunnel MTU : 1548                           Max Label Stack : 4
LSP Weight : 0
-----
Number of tunnel-table entries : 1
Number of tunnel-table entries with LFA : 0
=====
```

Weighted ECMP for SR-MPLS IGP shortcuts over SR-TE LSPs

The router supports weighted load balancing (weighted ECMP) for SR-OSPF packets forwarded over an ECMP set of SR-TE LSP IGP shortcuts. See the "Weighted load balancing IGP, BGP, and static route prefix packets over IGP shortcut" section in *7705 SAR Gen 2 Router Configuration Guide* for more information about weighted load balancing for IGP shortcut packets over SR-TE LSPs.

To enable weighted ECMP, configure the following commands.

```
configure router weighted-ecmp [strict]
configure router mpls lsp load-balancing-weight
```

If all the SR-TE LSPs in an ECMP set have a non-zero load-balancing weight, the router sprays the SR-OSPF packets in proportion to the **load-balancing-weight** normalized to a granularity of 64. If any LSPs have a zero load-balancing weight the router falls back to non-weighted ECMP, unless the **strict** option in the **weighted-ecmp** command is configured. If the **strict** option is configured, the router excludes LSPs with no (or zero) load-balancing weight from the ECMP set and performs weighted ECMP across the remaining LSPs.

If the next hop for a route resolves over SR-TE LSPs and direct IP interfaces so that both are included in the ECMP set, weighted ECMP can occur across all the SR-TE LSPs and IP interfaces. Otherwise, the router can use the relative metric of the SR-TE LSPs to exclude the interfaces from the set and the router only performs weighted ECMP across the SR-TE LSPs.

2.2.7 SR-TE LSP protection

The router supports local protection of a specific segment of an SR-TE LSP and end-to-end protection of the complete SR-TE LSP.

Whenever possible, an LFA next hop protects each path locally along the network. The protection of a node SID reuses the base LFA, TI-LFA, and remote LFA features used with SR shortest path tunnels. To augment the protection level, the SR OS adds the protection of an adjacency SID in the specific context of an SR-TE LSP. The user must enable the loopfree **[remote-lfa] [ti-lfa]** command in IS-IS or OSPF.

- **MD-CLI**

```
configure router isis loopfree-alternate remote-lfa
configure router isis loopfree-alternate ti-lfa
configure router ospf loopfree-alternate remote-lfa
configure router ospf loopfree-alternate ti-lfa
```

- **classic CLI**

```
configure router isis loopfree-alternates remote-lfa
configure router isis loopfree-alternates ti-lfa
configure router ospf loopfree-alternates remote-lfa
configure router ospf loopfree-alternates ti-lfa
```

An SR-TE LSP has state at the ingress LER only. The LSR has state for the node SIDs and adjacency SIDs that have labels programmed in the label stack of the received packet and that represent the part of the ERO of the SR-TE LSP on this router and downstream of this router. To provide protection for an SR-TE LSP, each LSR node must attempt to program a link-protect or node-protect LFA next hop in the ILM record of a node SID or of an adjacency SID, and the LER node must do the same in the LTN record of the SR-TE LSP. Details about this behavior are as follows:

- For an ILM record of a node SID of a downstream router that is not directly connected, the ILM of the node SID points to the backup NHLFE computed by the LFA SPF and programmed by the SR module for this node SID. Depending on the topology and LFA policy used, this can be a link-protect or node-protect LFA next hop.

This behavior is already supported in the SR shortest path tunnel feature at both the LER and the LSR. Consequently, an SR-TE LSP that transits at an LSR and matches the ILM of a downstream node SID automatically takes advantage of this protection, when enabled. If required, the user can disable node SID protection under the IGP instance by excluding the prefix of the node SID from the LFA.

- For an ILM record of a node SID of a directly connected router, the LFA SPF provides only link protection. The ILM or LTN record of this node SID points to the backup NHLFE of this LFA next hop. An SR-TE LSP that transits at an LSR and matches the ILM of a neighboring node SID automatically takes advantage of this protection, when enabled.



Note: Only link protection is possible in this case because packets matching this ILM record can either terminate on the neighboring router owning the node SID or can be forwarded to different next hops of the neighboring router (that is, to different next-next-hops of the LSR providing the protection). The LSR providing the connection does not have context to distinguish among all possible SR-TE LSPs and, therefore, can protect only the link to the neighboring router.

- For an ILM or LTN record of an adjacency SID, the handling is the same as for an ILM record of a node SID of a directly connected router.

When protecting an adjacency SID, the PLR first tries to select a parallel link to the node SID of the directly connected neighbor. The selection is based on the lowest interface ID. If no parallel links exist, regular LFA or remote LFA algorithms are applied to find a loopfree path to reach the node SID of the neighbor via other neighbors.

The ILM or LTN for the adjacency SID must point to this backup NHLFE and benefits from FRR link-protection. As a result, an SR-TE LSP that transits at an LSR and matches the ILM of a local adjacency SID automatically takes advantage of this protection, when enabled.

- For an ingress LER, the LTN record points to the SR-TE LSP NHLFE at the ingress LER, which itself points to the NHLFE of the SR shortest path tunnel to the node SID or adjacency SID of the first hop in the ERO of the SR-TE LSP. For this reason, the FRR link or node protection at an ingress LER is inherited directly from the SR shortest path tunnel.

When an adjacency to a neighbor fails, the following procedures are followed for both the LFA protected SID and the LFA unprotected SID of this adjacency in SR-MPLS. An adjacency can have both types of SIDs assigned by configuration. An LFA protected adjacency SID is eligible for LFA protection, but the following procedures apply even if a LFA backup was not programmed at the time of the failure. An LFA unprotected adjacency SID is not eligible for LFA protection.

- IGP withdraws the advertisement of the link TLV as well as its adjacency SID sub-TLV.
- The adjacency SID hold timer starts.
- The LTN and ILM records of the adjacency are kept in the datapath for as long as the adjacency SID hold timer is running. This allows packets to flow over the LFA backup path, when the adjacency is protected, and allows the ingress LER or PCE time to compute a new path of the SR-TE LSP after IGP converges.
- If the adjacency is restored while the adjacency SID hold timer is running, it remains programmed in the datapath with the retained SID values. However, the backup NHLFE may change if a new LFA

SPF runs while the adjacency SID hold timer is running. An update to the backup NHLFE is performed immediately following the LFA SPF. In all cases, the adjacency keeps its assigned SID label value.

- If the adjacency SID hold timer expires before the adjacency is restored, the SID is deprogrammed from the datapath and the label returned into the common pool where it was drawn from. Users of the adjacency (SR-TE LSP and SR Policy) are also informed. When the adjacency is subsequently restored, it gets assigned its allocated static-label value or a new dynamic-label value.
- A new PG-ID is assigned each time an adjacency comes back up. This PG-ID is used by the ILM and LTN of the adjacency SID and of all downstream node SIDs that resolve to a next hop over this adjacency.

The adjacency SID hold timer is configured using the **adj-sid-hold** command; it is activated when the adjacency to neighbor fails because of any of the following conditions.

- The network IP interface goes down because of a link or port failure or because the user performed a shutdown of the port.
- The user shuts down the network IP interface in the **configure router**, or **configure router ospf**, or **configure router isis** context.
- The adjacency SID hold timer is not activated if the user deletes an interface in the following contexts.

```
configure router ospf
configure router isis
```



Note: The adjacency SID hold timer does not apply to the ILM or LTN of a node SID, because NHLFE information is updated in the datapath as soon as IGP is converged locally and new primary and LFA backup next hops have been computed.

Although protection is enabled globally for all node SIDs and local adjacency SIDs when the user enables the **loopfree-alternate** option in IS-IS or OSPF at the LER and the LSR, applications may exist for which the user wants traffic to never divert from the strict hop computed by CSPF for an SR-TE LSP. In such cases, use the **sid-protection** command to disable protection for all adjacency SIDs formed over a specific network IP interface. Alternatively, configure a second unprotected SID for each adjacency using the **allocate-dual-sids** command.

The protection state of an adjacency SID is advertised in the B-flag of the IS-IS or OSPF adjacency SID sub-TLV.

2.2.7.1 Local protection

Whenever possible, an LFA next hop protects each path locally along the network. The protection of a SID node reuses the base LFA, TI-LFA, and remote LFA features introduced with segment routing shortest path tunnels. To augment the protection level, the SR OS adds the protection of an adjacency SID in the specific context of an SR-TE LSP. You must enable the loopfree **[remote-lfa] [ti-lfa]** command in IS-IS or OSPF.

- **MD-CLI**

```
configure router isis loopfree-alternate remote-lfa
configure router isis loopfree-alternate ti-lfa
configure router ospf loopfree-alternate remote-lfa
configure router ospf loopfree-alternate ti-lfa
```

- **classic CLI**

```
configure router isis loopfree-alternates remote-lfa
configure router isis loopfree-alternates ti-lfa
configure router ospf loopfree-alternates remote-lfa
configure router ospf loopfree-alternates ti-lfa
```

This behavior is already supported in the SR shortest path tunnel feature at both LER and LSR. An SR-TE LSP that transits at an LSR and that matches the ILM of a downstream SID node automatically takes advantage of this protection, when enabled. If required, SID node protection can be disabled under the IGP instance by excluding the prefix of the SID node from LFA.

2.2.7.2 End-to-end protection

This section provides a brief introduction to end to end protection for SR-TE LSPs.

End-to-end protection for SR-TE LSPs is provided using secondary or standby paths. Standby paths are permanently programmed in the datapath, whereas secondary paths are only programmed when they are activated. S-BFD is used to provide end-to-end connectivity checking. The **failure-action failover-or-down** command under the **bfd** context of the LSP is used to configure a switchover from the currently active path to an available standby or secondary path if the S-BFD session fails on the currently active path. If S-BFD is not configured, the router that is local to a segment can only detect failures of the top SID for that segment. End-to-end protection with S-BFD can be combined with local protection, but it is recommended that the S-BFD control packet timers be set to 1 second or more to allow sufficient time for any local protection action for a specific segment to complete without triggering S-BFD to go down on the end-to-end LSP path.

To prevent failure between the paths of an SR-TE LSP, that is to avoid, for example, a failure of a primary path that affects its standby backup path, then disjoint paths should be configured or the **srfg** command configured on the secondary paths.

As with RSVP-TE LSPs, SR-TE standby paths support the configuration of a path preference. This value is used to select the standby path to be used when more than one available path exists.

2.2.8 Static route resolution using SR-TE LSP

The user can forward packets of a static route to an indirect next hop over an SR-TE LSP programmed in the TTM by configuring the following static route tunnel binding command:

```
config>router>static-route-entry {ip-prefix/prefix-length} [mcast] indirect {ip-address}
  tunnel-next-hop
    - resolution {any | disabled | filter}
    - resolution-filter
      - [no] sr-te
        - [no] [lsp name1]
        - [no] [lsp name2]
        - .
        - .
        - [no] [lsp name-N]
      - exit
    - [no] disallow-igp
    - exit
  - exit
```

The user can select the **sr-te** tunnel type and either specify a list of SR-TE LSP names to use to reach the indirect next hop of this static route or allow the SR-TE LSPs to automatically select the indirect next hop in the TTM.

2.2.9 BGP shortcuts using SR-TE LSP

The user can forward packets of BGP prefixes over an SR-TE LSP programmed in TTM by configuring the following BGP shortcut tunnel binding command:

```
config>router>bgp>next-hop-resolution
- shortcut-tunnel
  - [no] family {ipv4}
    - resolution {any | disabled | filter}
    - resolution-filter
      - [no] sr-te
    - exit
  - exit
- exit
```

2.2.10 BGP labeled route resolution using SR-TE LSP

The user can enable SR-TE LSP, as programmed in TTM, for resolving the next hop of a BGP IPv4 or IPv6 (6PE) labeled route by enabling the following BGP transport tunnel command:

```
config>router>bgp>next-hop-res>
- labeled-routes
  - transport-tunnel
    - [no] family {label-ipv4 | label-ipv6 | vpn}
      - resolution {any | disabled | filter}
      - resolution-filter
        - [no] sr-te
      - exit
    - exit
  - exit
```

2.2.11 Service packet forwarding using SR-TE LSP

An SDP sub-type of the MPLS encapsulation type allows service binding to an SR-TE LSP programmed in the TTM by MPLS.

```
*A:Dut-A# configure service sdp 100 mpls create
-config>service>sdp$ sr-te-lsp lsp-name
```

The user can specify up to 16 SR-TE LSP names. The destination address of all LSPs must match that of the SDP far-end option. Service data packets are sprayed over the set of LSPs in the SDP using the same procedures as for tunnel selection in ECMP. However, each SR-TE LSP can have up to 32 next hops at the ingress LER when the first segment is a node SID-based SR tunnel. Consequently, service data packets are forwarded over one of a maximum of 16×32 next hops. The **tunnel-far-end** option is not supported. In addition, the **mixed-lsp-mode** option does not support the **sr-te** tunnel type.

The signaling protocol for the service labels of an SDP that is using an SR-TE LSP can be configured to static (**off**), T-LDP (**tldp**), or BGP (**bgp**).

Use the following command syntax to configure an SR-TE LSP used in VPRN auto-bind.

```
config>service>vprn>
  - auto-bind-tunnel
    - resolution {any | disabled | filter}
    - resolution-filter
      - [no] sr-te
    - exit
  - exit
```

Both VPN-IPv4 and VPN-IPv6 (6VPE) are supported in a VPRN service using segment routing transport tunnels with the **auto-bind-tunnel** command.

Use the following command syntax with the BGP EVPN service.

```
config>service>vpls>bgp-evpn>mpls>
  - auto-bind-tunnel
    - resolution {any | disabled | filter}
    - resolution-filter
      - [no] sr-te
    - exit
  - exit
```

The following service contexts are supported with SR-TE LSP:

- VLL, LDP VPLS, IES/VPRN spoke-interface, R-VPLS, BGP EVPN
- BGP-AD VPLS, BGP-VPLS, BGP VPWS when the **use-provisioned-sdp** option is enabled in the binding to the PW template
- intra-AS BGP VPRN for VPN-IPv4 and VPN-IPv6 prefixes with both auto-bind and explicit SDP
- inter-AS options B and C for VPN-IPv4 and VPN-IPv6 VPRN prefix resolution
- IPv4 BGP shortcut and IPv4 BGP labeled route resolution
- IPv4 static route resolution
- multicast over IES/VPRN spoke interface with **spoke SDP** riding a SR-TE LSP

2.2.12 Datapath support

To support SR-TE in the datapath, the ingress LER must push a label stack where each label represents a hop, a TE link, or a node, in the ERO for the LSP path computed by the router or the PCE. However, only the label and the outgoing interface to the first strict or loose hop in the ERO factor into the forwarding decision of the ingress LER, because the SR-TE LSP only needs to track the reachability of the first strict or loose hop. This represents the NHLFE of the SR shortest path tunnel to the first strict or loose hop.

To ensure that its NHLFE is readily available, the SR OS stores the SR shortest path tunnel to a downstream node SID or adjacency SID in the tunnel table. The rest of the label stack is not meaningful to the forwarding decision. In this guide, "super NHLFE" refers specifically to this part of the label stack because it can have a much larger size.

An SR-TE LSP is modeled in the ingress LER datapath as a hierarchical LSP, with the super NHLFE tunneled over the NHLFE of the SR shortest path tunnel to the first strict or loose hop in the SR-TE LSP path ERO.

The following are characteristics of this model.

- The model saves on NHLFE usage. When many SR-TE LSPs travel to the same first hop, they ride the same SR shortest path tunnel, and each consumes one super NHLFE. However, the SR-TE LSPs point to a single NHLFE or set of NHLFEs if ECMP exists for the first strict or loose hop of the first-hop SR tunnel.

In addition, the ingress LER does not need to program a separate backup super NHLFE. Instead, the single super NHLFE automatically begins forwarding packets over the LFA backup path of the SR tunnel to the first hop as soon as it is activated.

- When the path of a SR-TE LSP contains a maximum of two SIDs, that is the destination SID and one additional loose or strict-hop SID, the SR-TE LSP uses a hierarchy consisting of a regular NHLFE pointing to the NHLFE of top SID corresponding to the first loose or strict hop.
- If the first segment is a node SID tunnel and multiple next hops exist, ECMP spraying is supported at the ingress LER.
- If the first-hop SR tunnel, node, or adjacency SID goes down, the SR module informs MPLS that the outer tunnel is down, and MPLS brings the SR-TE LSP down and requests SR to delete the SR-TE LSP in the IOM.

The datapath behavior at the LSR and the egress LER for an SR-TE LSP is similar to that of a shortest path tunnel, because there is no tunnel state in these nodes. Packet forwarding is based on processing the incoming label stack, consisting of a node SID or adjacency SID label, or both. If the ILM is for a node SID and multiple next hops exist, ECMP spraying is supported at the LSR.

The link-protect LFA backup next hop for an adjacency SID can be programmed at the ingress LER and LSR nodes. See [SR-TE LSP protection](#) for more information.

A maximum of 12 labels, including all transport, service, hash, and OAM labels, can be pushed. The label stack size for the SR-TE LSP can be 1 to 11 labels, with a default value of 6.

The maximum value of 11 is obtained for an SR-TE LSP whose path is not protected via FRR backup and with no entropy or hash label feature enabled when such an LSP is used as a shortcut for an IGP IPv4/IPv6 prefix or as a shortcut for BGP IPv4/IPv6. In this case, the IPv6 prefix requires pushing the IPv6 explicit-null label at the bottom of the stack. This leaves 11 labels for the SR-TE LSP.

The default value of 6 is obtained in the worst cases, such as forwarding a vprn-ping packet for an inter-AS VPN-IP prefix in Option C:

6 SR-TE labels + 1 remote LFA SR label + BGP 8277 label + ELI (RFC 6790) + EL (entropy label) + service label + OAM Router Alert label = 12 labels.

The label stack size manipulation includes the following LER and LSR roles:

LER role

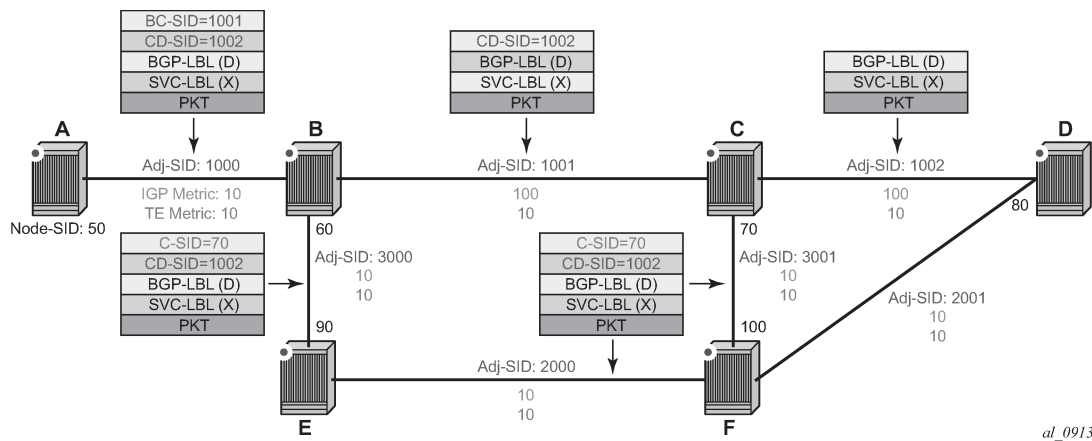
- push up to 12 labels
- pop-up to 8 labels of which 4 labels can be transport labels

LSR role

- pop-up to 5 labels and swap one label for a total of 6 labels
- LSR hash of a packet with up to 16 labels

The following figure shows an example of the label stack pushed by the ingress LER and by an LSR acting as a PLR.

Figure 21: SR-TE LSP label stack programming



On node A, the user configures an SR-TE LSP to node D with a list of explicit strict hops mapping to the adjacency SID of links A-B, B-C, and C-D.

Ingress LER A programs a super NHLFE consisting of the label for the adjacency over link C-D and points the NHLFE to the already programmed NHLFE of the SR tunnel of its local adjacency over link A-B. The latter NHLFE has the top label and also the outgoing interface to send the packet to.



Note: SR-TE LSP does not consume a separate backup super NHLFE; it only points the single super NHLFE to the NHLFE of the SR shortest path tunnel it is riding. When the latter activates its backup NHLFE, the SR-TE LSP automatically forwards over it.

LSR node B has already programmed the primary NHLFE for the adjacency SID over link C-D and has the ILM with label 1001 point to it. In addition, node B preprograms the link-protect LFA backup next hop for link B-C and points the same ILM to it.



Note: There is no super NHLFE at node B, because it only deals with programming the ILM and primary or backup NHLFE of its adjacency SIDs and its local and remote node SIDs.

VPN service in node A forwards a packet to the VPN-IPv4 prefix X advertised by BGP peer D. [Figure 21: SR-TE LSP label stack programming](#) shows the resulting datapath at each node for the primary path and for the FRR backup path at LSR B.

2.2.12.1 SR-TE LSP metric and MTU settings

The MPLS module assigns the maximum LSP metric value (16777215) to a TE LSP when the local router provides the hop-to-label translation for its path. If a TE LSP uses the local CSPF (**path-computation-method local-cspf** option enabled) or PCE for path computation (**path-computation-method pce** option enabled) or delegates control to the PCE (**pce-control** enabled), the PCE returns the computed LSP IGP or TE metric in the PCReq and PCUpd messages.

In all cases, using the **config router mpls lsp metric** command to configure an admin metric overrides the returned value.

The MTU computations are as follows:

- The MTU setting of an SR-TE LSP is derived from the MTU of the outgoing SR shortest path tunnel it is riding, adjusted with the size of the super NHLFE label stack size.

The following calculation is used:

$$\text{SR_Tunnel_MTU} = \text{MIN} \{ \text{Cfg_SR_MTU}, \text{IGP_Tunnel_MTU} - (1 + \text{frr-overhead}) \times 4 \}$$

where:

- Cfg_SR_MTU is the MTU configured for all SR tunnels within a specific IGP instance using the **config router ospf segment-routing tunnel-mtu** or **config router isis segment-routing tunnel-mtu** command. If no value is configured, the SR tunnel MTU is fully determined by the IGP interface calculation, described in the next bullet point.
- IGP_Tunnel_MTU is the minimum of the IS-IS or OSPF interface MTU among all the ECMP paths or among the primary and LFA backup paths of the SR tunnel.
- *frr-overhead* is set to:
 - value of **ti-lfa [max-sr-frr-labels labels]** if **loopfree-alternates** and **ti-lfa** are enabled in this IGP instance
 - 1 if **loopfree-alternates** and **remote-lfa** are enabled but **ti-lfa** is disabled in this IGP instance
 - 0 for all other cases

This calculation is performed by IGP and passed to the SR module each time there is a change because of an updated resolution of the node SID.

The SR OS also provides the MTU for the adjacency SID tunnel because it is needed in an SR-TE LSP if the first hop in the ERO is an adjacency SID. In this case, the calculation for SR_Tunnel_MTU, initially introduced for a node SID tunnel, is applied to derive the MTU of the adjacency SID tunnel.

- The MTU of the SR-TE LSP is derived as follows:

$$\text{SRTE_LSP_MTU} = \text{SR_Tunnel_MTU} - \text{numLabels} \times 4$$

where:

- SR_Tunnel_MTU is the MTU SR tunnel shortest path that the SR-TE LSP is riding. The formula for this is provided in the previous bullet point.
- numLabels is the number of labels found in the super NHLFE of the SR-TE LSP. At the LER, the super NHLFE points to the SR tunnel NHLFE, which itself has a primary and a backup NHLFE.

This calculation is performed by the SR module and is updated each time the SR-TE LSP path changes or the SR tunnel it is riding is updated.



Note: The above calculated SR-TE LSP MTU is used for the determination of an SDP MTU and for checking the Layer 2 service MTU. In the case of fragmentation of IP packets forwarded in GRT or in a VPRN over an SR-TE LSP, the IOM always deducts the worst-case MTU (12 labels) from the outgoing interface MTU when deciding whether to fragment the packet. In this case, the preceding formula is not used.

2.2.12.2 LSR hashing on SR-TE LSPs

The LSR supports hashing up to a maximum of 16 labels in a stack. The LSR is able to hash on the IP headers when the payload below the label stack is IPv4 or IPv6, including when a MAC header precedes it (**eth-encap-ip** command option). Alternatively, it is able to hash based only on the labels in the stack, which may include the entropy label (EL) or the hash label. See the *7705 SAR Gen 2 MPLS Guide* for more information about the hash label and entropy label features.

When the **hash-label** command option is enabled in a service context, a hash label is always inserted at the bottom of the stack as per RFC 6391.

The EL feature, as specified in RFC 6790, indicates the presence of a flow on an LSP that should not be reordered during load balancing. It can be used by an LSR as input to the hash algorithm. The ELI is used to indicate the presence of the EL in the label stack. The ELI, followed by the actual EL, is inserted immediately below the transport label for which the EL feature is enabled. If multiple transport tunnels have the EL feature enabled, the ELI and EL are inserted below the lowest transport label in the stack.

The EL feature is supported with an SR-TE LSP. See the *7705 SAR Gen 2 MPLS Guide* for more information.

The LSR hashing operates as follows:

- If the **lbl-only** hashing command option is enabled, or if one of the other LSR hashing options is enabled but an IPv4 or IPv6 header is not detected below the bottom of the label stack, the LSR parses the label stack and hashes only on the EL or hash label.
- If the **lbl-ip** command option is enabled, the LSR parses the label stack and hashes on the EL or hash label and the IP headers.
- If the **ip-only** or **eth-encap-ip** command option is enabled, the LSR hashes on the IP headers only.

2.2.13 SR-TE Auto-LSP

The SR-TE auto-LSP feature supports auto-creation of the following types of LSPs:

- SR-TE mesh
- SR-TE one-hop
- SR-TE on-demand

The SR-TE mesh LSP feature binds an SR-TE mesh P2P LSP template with one or more prefix lists. When the TE database discovers a router that has an ID matching an entry in the prefix list, the database triggers MPLS to instantiate an SR-TE LSP to the router using the LSP parameters in the LSP template.

The SR-TE one-hop LSP feature activates an SR-TE one-hop P2P LSP template. In this case, the TE database tracks each TE link that is made to a directly connected IGP neighbor. The database then instructs MPLS to instantiate an SR-TE LSP with the following parameters:

- the source address of the local router
- an outgoing interface matching the interface index of the TE link
- a destination address matching the router ID of the neighbor on the TE link

In both these types of SR-TE auto-LSP, the router hop-to-label translation or local CSPF computes the label stack required to instantiate the LSP path.

The SR-TE on-demand LSP feature creates an LSP using an SR-TE on-demand P2P LSP template. When an imported BGP route matches an entry in a policy statement with an MPLS create tunnel action, an LSP is created to the next hop for the route. If a route admin tag policy is applied when the route is imported, only an auto-LSP with a template containing a matching **admin-tag** is created. The SR-TE on-demand LSP supports path computation using hop-to-label translation, local-CSPF, or a PCE.



Note: An SR-TE mesh or one-hop auto-LSP can be reported to a PCE but cannot be delegated or have its paths computed by PCE. An SR-TE on-demand LSP can also be controlled and have its path computed by a PCE, as well as being reported to a PCE.

2.2.13.1 Feature configuration

About this task

This feature uses three SR-TE LSP template types: one-hop P2P, on-demand P2P, and mesh P2P. For the one-hop P2P and mesh P2P types, the configuration of the commands is the same as that of the RSVP-TE auto-LSP.

Procedure

Step 1. Create an LSP template using one of the following commands, depending on the type of auto-LSP required.

MD-CLI

```
configure router mpls lsp-template type p2p-sr-te-mesh
configure router mpls lsp-template type p2p-sr-te-one-hop
configure router mpls lsp-template type p2p-sr-te-on-demand
```

classic CLI

```
configure router mpls lsp-template mesh-p2p-srte
configure router mpls lsp-template one-hop-p2p-srte
configure router mpls lsp-template on-demand-p2p-srte
```

Step 2. In the template, configure the common LSP- and path-level parameters or options shared by all LSPs using this template.



Note: These LSP template types contain the SR-TE LSP-specific commands and other LSP or path commands that are common to RSVP-TE and SR-TE LSPs and are supported by the existing RSVP-TE LSP template.

Step 3. Bind the LSP templates as follows:

- For the SR-TE mesh P2P LSP template, use the **configure router mpls lsp-template policy** command.
- For the SR-TE one-hop P2P LSP template, use the **configure router mpls lsp-template one-hop** command.
- For the on-demand SR-TE LSP template, bind the template to the creation of SR-TE auto-LSPs using the **configure router mpls auto-lsp** command and configure the **create-mpls-tunnel** command as an action in a route import policy statement.

See [Configuring and operating SR-TE](#) for an example of SR-TE auto-LSP creation using an LSP template type **mesh-p2p-srte**.

2.2.13.2 Automatic creation of an SR-TE mesh LSP

The **auto-lsp** command binds an LSP template of type **mesh-p2p-srte** with one or more prefix lists. When the TE database discovers a router that has a router ID matching an entry in the prefix list, it triggers MPLS to instantiate an SR-TE LSP to that router using the LSP parameters in the LSP template.

The prefix match can be exact or longest. Prefixes in the prefix list that do not correspond to a router ID of a destination node cannot match.

The path of the LSP is that of the default path name specified in the LSP template. The hop-to-label translation tool or the local CSPF determines the node SID and adjacency SID corresponding to each loose and strict hop in the default path definition, respectively.

The LSP has an auto-generated name using the following structure:

TemplateName-DestIpv4Address-TunnelId

where:

- *TemplateName* is the name of the template
- *DestIpv4Address* is the address of the destination of the auto-created LSP
- *TunnelId* is the TTM tunnel ID

In SR OS, an SR-TE LSP uses three different identifiers:

- LSP Index is used for indexing the LSP in the MIB table shared with RSVP-TE LSP. The LSP Index range is as follows:
 - provisioned SR-TE LSP: 65536 to 81920
 - SR-TE auto-LSP: 81921 to 131070
- LSP Tunnel ID is used in the interaction with the PCC or PCE. Range: 1 to 65536
- TTM Tunnel ID is the tunnel ID service, shortcut, and steering applications use to bind to the LSP. Range: 655362 to 720897

The path name is the default path specified in the LSP template.



Note: This feature is limited to SR-TE LSPs that are controlled by the router (PCC-controlled), where the path is provided using the hop-to-label translation or the local CSPF path computation method.

2.2.13.3 Automatic creation of an SR-TE one-hop LSP

Although the provisioning model and CLI syntax differ from that of a mesh LSP by the absence of a prefix list, the actual behavior is quite different. When the **one-hop-p2p** command is executed, the TE database keeps track of each TE link that comes up to a directly connected IGP neighbor. It then instructs the MPLS to instantiate an SR-TE LSP with the following parameters:

- the source address of the local router
- an outgoing interface matching the interface index of the TE link
- a destination address matching the router ID of the neighbor on the TE link

The hop-to-label translation or the local CSPF returns the SID for the adjacency to the remote address of the neighbor on this link. Therefore, the **auto-lsp** command binding an LSP template of type **one-hop-p2p-srte** with the **one-hop** option results in one SR-TE LSP instantiated to the IGP neighbor for each adjacency over any interface.

Because the local router installs the adjacency SID to a link regardless of whether the neighbor is SR-capable, the TE-DB finds the adjacency SID and a one-hop SR-TE LSP can still come up to such a neighbor. However, remote LFA using the neighbor's node SID does not protect the adjacency SID or the one-hop SR-TE LSP because the node SID is not advertised by the neighbor.

The LSP has an auto-generated name using the following structure:

TemplateName-DestIpv4Address-TunnelId

where:

- *TemplateName* = the name of the template
- *DestIpv4Address* = the address of the destination of the auto-created LSP
- *TunnelId* = the TTM tunnel ID

The path name is the default path specified in the LSP template.



Note: This feature is limited to an SR-TE LSP that is controlled by the router (PCC-controlled) and the path labels are provided by the hop-to-label translation or the local CSPF path computation method.

2.2.13.4 Automatic creation of an on-demand SR-TE LSP

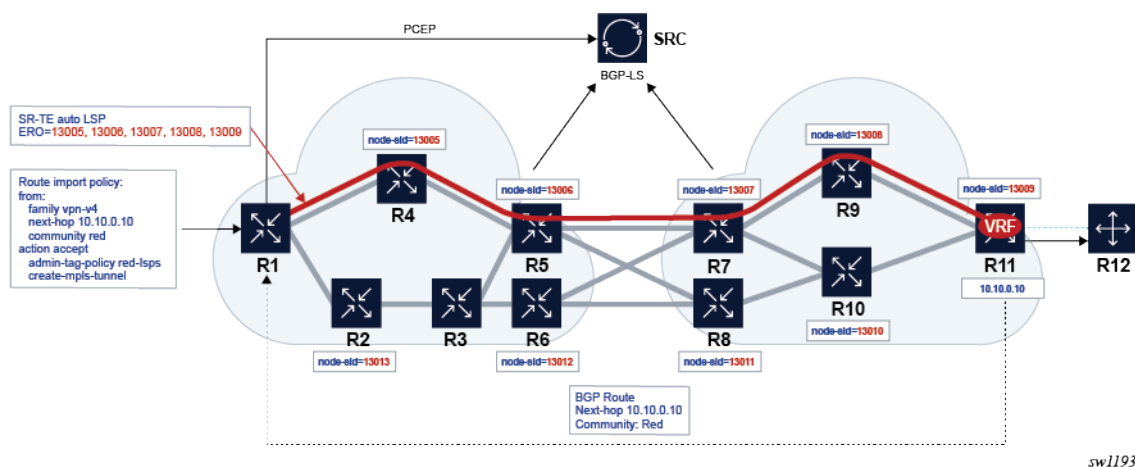
The SR-TE on-demand LSP simplifies provisioning for networks that may or may not be managed by a network service manager, such as the Nokia NSP. Instead of using a full mesh, LSPs can be automatically created on-demand when a suitable tunnel does not exist for a specific BGP prefix next hop. The prefix could be for VPRN, EVPN, BGP-LU, or BGP shortcut routes. Both intra-domain and inter-domain use cases are supported.

This mechanism is an extension of the LSP **admin-tag** and auto-LSP mechanisms and applies to the following objects:

- VPRN auto-bind-tunnel
- EVPN VPWS auto-bind-tunnel
- EVPN VPLS auto-bind-tunnel
- BGP-LU, both as an LER and LSR at an ABR or ASBR
- BGP shortcuts

The following figure shows an application of SR-TE on-demand LSPs.

Figure 22: VPRN example of an on-demand SR-TE LSP



This example combines route transport coloring and auto LSPs to simplify provisioning for intent-based networking for specific services. In this use case, intent means the ability to meet traffic engineering

requirements for a service. This could be, for example, a delay or loss, or the ability to steer the service traffic to avoid LSPs that transit specific geographies or prefer those that take another route.

In this example, a BGP route is advertised for a VRF in a PE for a VPRN service. An extended color community is assigned to the route. This color implies an intent associated with the transport requirements.

When this route is imported at the head-end PE, the router performs the following steps:

1. The route is matched in a route-import policy.
2. An **admin-tag** policy called red-lsps is applied.
3. A trigger action occurs to create an MPLS tunnel to the BGP next-hop for the route.

This causes the head-end router to create an SR-TE auto-LSP that matches the red-lsps **admin-tag** policy and steers the traffic associated with "red" routes to the far-end PE, into the red LSP. This SR-TE auto-LSP is created based on the configuration in the matching LSP template.

SR OS also offers the ability to use the local CSPF, hop-to-label-translation, or a PCE to provide a path for the red LSP. This is determined by a configuration in the matching LSP template.

2.2.13.4.1 Deletion of on-demand SR-TE LSPs

SR-TE on-demand P2P auto LSPs are removed by the router in all the following cases.

- The classic CLI **no auto-lsp** (or MD-CLI **delete auto-lsp**) command is executed. This triggers MPLS to remove auto-LSPs created by this command.
- The **no create-mpls-tunnel** is configured in a policy statement that previously had **create-mpls-tunnel** configured. This triggers a reevaluation of the policy statement and potentially triggers BGP to inform MPLS that it no longer needs a tunnel.
- BGP tracks the binding of a route to an **admin-tag-policy**. If an **admin-tag-policy** name in a policy statement action changes, the policy is reevaluated, which could change the binding. This may result in a request to create a new tunnel or delete an existing tunnel. However, if the contents of an **admin-tag-policy** that is referenced in a policy statement action change, BGP does not react (for example, request the creation or deletion of a tunnel), although a subsequent route resolution may change.
- MPLS reacts to **admin-tag** changes in the LSP template. When this occurs, it reevaluates the **admin-tag-policy** associated with a request from BGP and deletes or creates tunnels accordingly.
- If a new LSP is created that is not an on-demand LSP and is preferred to an existing on-demand LSP, BGP can resolve the next hop over the new LSP and traffic moves to it. In this case, the system does not remove the older less-preferred auto-LSP, which was created through an on-demand LSP trigger, until the next hops are removed.
- If the LSP template is shut down, all associated LSPs are administratively disabled. To delete the LSP template you must first shut it down, using a **no auto-lsp** command in classic CLI or **delete auto-lsp** command in MD-CLI. This removes all the auto-LSPs that are using the template.

2.2.13.4.2 Configuring SR-TE on-demand LSPs

About this task

Configure SR-TE on-demand LSPs using the steps in this section.

Procedure

Step 1. Define a policy statement to import the route, as shown in the following example:

Example

```
configure>router>policy-options>policy-statement
    entry
    from
        family <family>
        next-hop <ip-address>
        community <comm-name>
    action accept
        admin-tag-policy <admin-tag-policy-name>
        create-mpls-tunnel
```

Step 2. Configure the auto-LSP under MPLS with the template type **on-demand-p2p-srte**.

The **create-mpls-tunnel** action is supported for the following address families:

- vpn-ipv4
- vpn-ipv6
- evpn
- label-ipv4
- label-ipv6
- ipv4
- ipv6

The router-policy action assigns an **admin-tag-policy** to the routes that are imported with a specific next hop and match a specified extended community. In most applications, the extended community is the transport color extended community. The **create-mpls-tunnel** command action causes BGP to send the next hop and the include and exclude constraints in the **admin-tag-policy** (if one was assigned to a route by the policy statement) to the MPLS application.

When such a policy statement is applied in the context of a specific VRF, the **create-mpls-tunnel** command trigger is only actioned by BGP on a per-next-hop basis.

This type of LSP template supports PCE computation, control, and the fallback path computation method if the PCE is unreachable. The auto-LSP is configured using the following command:

```
configure>router>mpls
    auto-lsp <on-demand-p2p-srte-template-name>
```

The LSP template may contain an LSP **admin-tag-policy**. MPLS takes the next hop, and the **admin-tag** command includes or excludes constraints from BGP and matches them against the auto-LSP statement with a template with an **admin-tag** command that conforms to the **admin-tag-policy** constraints.

If BGP does not pass any **admin-tag-policy** constraints, MPLS only matches against LSP templates that do not have the **admin-tag** command configured.

If the next-hop and **admin-tag-policy** match more than one auto-LSP statement, an LSP is created for each matching entry. This results in an ECMP set to the next hop.



Note: Each LSP may have a different **admin-tag** value, but it is an ECMP next-hop tunnel from the perspective of the colored route that triggers the tunnel creation.

A new SR-TE LSP is consequently created to the next hop passed by BGP according to the parameters contained in the LSP template.

The router tracks the binding between BGP triggers and on-demand LSPs that are successfully created and deleted toward a specified BGP next-hop matching an **admin-tag-policy**.

2.2.13.5 Forwarding contexts supported with SR-TE auto-LSP

The following are the forwarding contexts that can be used by an auto-LSP:

- resolution of IPv4 BGP labeled routes and IPv6 BGP labeled routes (6PE) in the TTM
- resolution of IPv4 BGP route in the TTM (BGP shortcut)
- resolution of IPv4 static route to indirect next hop in the TTM
- VPRN and BGP-EVPN auto-bind for both IPv4 and IPv6 prefixes

The auto-LSP cannot be used in a provisioned SDP for explicit binding by services. Therefore, an auto-LSP can also not be used directly for auto-binding of a PW template with the **use-provisioned-sdp** option in BGP-AD VPLS or FEC129 VLL service. However, an auto-binding of a PW template to an LDP LSP, which is then tunneled over an SR-TE auto-LSP is supported.

2.2.14 Allocation and binding of labels to SR-TE LSPs

SR OS supports the allocation and binding of labels to SR-TE LSPs. The LSPs must be named LSPs.

The result of a binding SID label is the programming of an ILM with a swap operation pointing to the LSP NHLFE.

A single binding SID label can be allocated to a specific LSP.

Named LSPs

Use the commands in the following context to configure named LSPs.

```
configure router mpls lsp
```

Use the following command to configure the binding SID label value for named LSPs.

```
configure router mpls lsp binding-sid
```

The value of the binding SID label must be within the label block that is reserved for binding SID labels. The reserved label block is configured like any other reserved block. Use the following command to reference the reserved label block for statically configured binding SIDs.

```
configure router mpls lsp-bsid-block
```

A binding SID label can be assigned or removed at any time. The node that allocates the label is considered to be the owner.

2.2.15 SR-TE LSP traffic statistics

The collection of traffic statistics on SR-TE LSPs using either a named LSP or SR-TE templates is available on egress or ingress LER. Also, traffic statistics cannot be recorded into an accounting file.

SR-TE LSP statistics are provided without any forwarding class or QoS profile distinction. However, traffic statistics are recorded and made available for each path of the LSP (primary and backup). Statistic indexes are only allocated at the time the path is effectively programmed, are maintained across switchover for primary and standby LSPs only, and are released if egress statistics are disabled or the LSP is deleted.

2.2.15.1 Rate statistics

About this task

SR OS also provides traffic rate statistics.

The frequency at which rate statistics are determined is configured in the accounting policy using the **collection-interval** command. The minimum interval is 5 minutes.

Rate statistics for SR-TE LSPs cannot be written to an accounting file. The **to no-file** command must be configured in the accounting policy.

Rate statistics are provided in pkt/s and Mb/s. Rate statistics are provided as an aggregate across all paths of the LSP for which a statistical index has been assigned, and for all forwarding classes in- or out-of-profile.

Rate statistics are only available on egress of the ingress LER. At least two samples are needed to determine a rate.

For SR-TE LSPs, including template-based LSPs, the user enables this capability by performing the following tasks:

Procedure

- Step 1.** Configure an accounting policy that uses the record **combined-mpls-srte-egress**.
- Step 2.** Assign the configured accounting policy to a specific LSP (or template).
- Step 3.** Enable stats collection.

2.2.16 SR-TE label stack checks

This section describes the SR-TE label stack checks.

2.2.16.1 SR-TE label stack check for services and shortcuts

If a packet forwarded in a service results in the net label stack size being pushed on the packet to exceed the maximum label stack supported by the router, the packet is dropped on the egress. Each service and shortcut application on the router performs a check of the resulting net label stack after pushing all the labels required for forwarding the packet in that context.

To that effect, the MPLS module populates each SR-TE LSP in the TTM with the maximum transport label stack size, which consists of the sum of the values in **max-sr-labels** *label-stack-size* and **additional-frr-labels** *labels*.

Each service or shortcut application then adds the additional, context-specific labels such as service label, entropy or hash label, and control-word, required to forward the packet in that context and check that the resulting net label stack size does not exceed the maximum label stack supported by the router.

If the check succeeds, the service binds or the prefix resolves to the SR-TE LSP.

If the check fails, the service does not bind to this SR-TE LSP. Instead, the service either finds another SR-TE LSP or another tunnel type to bind to, if the user has configured other tunnel types. Otherwise, the service goes down. If the service uses an SDP with one or more SR-TE LSP names, the spoke SDP bound to this SDP remains operationally down as long as at least one SR-TE LSP fails the check. In this case, a new spoke SDP flag is displayed in the show output of the service: `labelStackLimitExceeded`. Similarly, the prefix is not resolved to the SR-TE LSP and is either resolved to another SR-TE LSP or another tunnel type, or becomes unresolved.

The value of **additional-frr-labels labels** is checked against the maximum value across all IGP instances of the *frr-overhead* parameter. This parameter is computed within a specific IGP instance. The following table lists the parameter values.

Table 8: Values of the *frr-overhead* parameter

Condition	frr-overhead parameter value
segment-routing is disabled in the IGP instance	0
segment-routing is enabled but remote-lfa is disabled	0
segment-routing is enabled and remote-lfa is enabled	1

If the user configures or changes the configuration of the **additional-frr-labels** command, MPLS ensures that the new value accommodates the *frr-overhead* value across all IGP instances. Consider the sequence in the following example:

1. The user configures the **config router isis loopfree-alternate remote-lfa** command.
2. The user creates an SR-TE LSP or changes the configuration of an existing one, as follows:

mpls>lsp>max-sr-labels 10 additional-frr-labels 0.



Note: Performing a **no shutdown** of the new LSP or changing the existing LSP configuration is blocked because the IS-IS instance enabled remote LFA, which requires one additional label on top of the 10 SR labels of the primary path of the SR-TE LSP.

If the check is successful, MPLS adds **max-sr-labels** and **additional-frr-labels** and checks that the value is less than or equal to the maximum label stack supported by the router. MPLS then populates the value of {**max-sr-labels** + **additional-frr-labels**}, along with tunnel information in TTM, and also passes **max-sr-labels** to the PCEP module.

Conversely, if the user attempts a configuration change that may result in a change to the computed *frr-overhead* value, IGP checks that all SR-TE LSPs can properly account for the overhead; if the check fails, the change is rejected. On the IGP, enabling the **remote-lfa** command may cause the *frr-overhead* value to change. Consider the sequence in the following example:

1. An MPLS LSP is administratively enabled and has **mpls>lsp>max-sr-labels 10 additional-frr-overhead 0** configured.
2. The current configuration in IS-IS has the **loopfree-alternate** command disabled.

3. The user attempts to configure **isis loopfree-alternate remote-lfa**. This changes the *frr-overhead* value to 1.

This configuration change is blocked.

2.2.16.2 Control plane handling of egress label stack limitations

As described in [Datapath support](#), the egress IOM can push a maximum of 12 labels; however, this number may be reduced if other fields are pushed into the packets. For example, for a VPRN service, the ingress LER can send an IP VPN packet with 12 labels in the stack, including one service label, one label for OAM, and 10 transport labels. However, if entropy is configured, the number of transport labels is reduced by two (EL and ELI). Similarly, for EVPN services, the egress IOM may push specific fields that reduce the total number of supported transport labels.

To avoid silent packet drops in cases where the egress IOM cannot push the required number of labels, SR OS implements a set of procedures that prevent the system from sending packets if it is determined that the SR-TE label stack to be pushed exceeds the number of bytes that the egress IOM can put on the wire.

The following table describes the label stack egress IOM restrictions for IP-VPN and EVPN services for the 7705 SAR Gen 2.



Note: The 7705 SAR Gen 2 does not support ESI labels (EVPN multi-homing). ESI label information is included in this section for reference only.

Table 9: Label stack egress IOM restrictions on FP-based hardware for IPVPN and EVPN services

Features that reduce the label stack		Source service type			
		IP-VPN (VPRN)	EVPN-IFL (VPRN)	EVPN VPLS or EVPN Epipe	EVPN-IFF (R-VPLS)
Always Computed ³	Service Label	1	1	1	1
	OAM Label	1	1	0	0
	Control Word (CW)	0	0	1	1
Computed if configured ⁴	Hash Label (mutually exclusive with EL)	1	0	0	0
	EL+ELI	2	2	2	2

³ These rows indicate the number of labels that the system assumes are always used on a specific service. For example, the system always computes two labels to be reduced from the total number of labels for VPRN services with EVPN-IFL (EVPN Interface-less model enabled).

⁴ These rows indicate the number of labels that the system subtracts from the total only if they are configured on the service. For example, on VPRN services with EVPN-IFL, if the user configures hash-label, the system computes one additional label. If the user configures entropy-label, the system deducts two labels instead.

Features that reduce the label stack	Source service type			
	IP-VPN (VPRN)	EVPN-IFL (VPRN)	EVPN VPLS or EVPN Epipe	EVPN-IFF (R-VPLS)
Required Labels ⁵	2	2	2	2
Required Labels + Options ⁵	4	4	4	4
Maximum available labels ⁶	12	12	10	9
Maximum available transport labels without options ⁷	10	10	8	7
Maximum available transport labels with options ⁷	8	8	6	5

The total number of labels configured in the **max-sr-labels label-stack-size [additional-frr-labels labels]** command must not exceed the labels indicated in the "Maximum available transport labels without options" and "Maximum available transport labels with options" rows in the preceding table. If the configured LSP labels exceed the available labels listed in the preceding table, the BGP route next hop for the LSP is not resolved and the system does not even try to send packets to that LSP.

For example, for a VPRN service with EVPN-IFL where the user configures EL, the maximum available transport labels is eight. If an IP Prefix route for next-hop X is received for the service and the SR-TE LSP to-X is the best tunnel to reach X, the system checks that (**max-sr-labels + additional-frr-labels**) is less than or equal to eight. Otherwise, the IP Prefix route is not resolved.

The same control plane check is performed for other service types, including IP shortcuts, spoke SDPs on IP interfaces, spoke SDPs on Epipes, VPLS, and R-VPLS. In all cases, the spoke SDP is brought down if the configured (**max-sr-labels + additional-frr-labels**) is greater than the maximum available transport labels. The following table indicates the maximum available transport labels for IP shortcuts and spoke SDP services on the 7705 SAR Gen 2.

⁵ These rows indicate the number of labels that the system deducts from the total number.

⁶ This row indicates a different number depending on the service type and the inner encapsulation used by each service, which reduces the maximum number of labels to push on egress. For example, while the number of labels for VPRN services is 12, the maximum number for VPLS and Epipe services is 10 (to account for space for an inner Ethernet header).

⁷ This row indicates the maximum SR-TE labels that the system can push when sending service packets on the wire.

Table 10: Maximum available transport labels for IP shortcuts and spoke SDP services

Features that reduce the label stack		Source service type				
		IP shortcuts	Spoke-SDP interface	Spoke-SDP Epipe	Spoke-SDP VPLS	Spoke-SDP R-VPLS
Always Computed ⁸	Service Label	0	1	1	1	1
	OAM Label	0	1	1	1	0
	IPv6 label	1	0	0	0	0
Computed if configured ⁹	Hash Label (mutually exclusive with EL)	0	1	1	1	1
	EL+ELI	2	2	2	2	2
	CW	0	1	1	1	1
Required Labels ¹⁰		1	2	2	2	1
Required Labels + Options ¹⁰		3	5	5	5	4
Maximum available labels ¹¹		12	9	10	10	8
Maximum available transport labels without options ¹²		11	7	8	8	7
Maximum available transport labels with options ¹²		9	4	5	5	4

In general, the labels shown in [Table 9: Label stack egress IOM restrictions on FP-based hardware for IPVPN and EVPN services](#) and [Table 10: Maximum available transport labels for IP shortcuts and spoke SDP services](#) are valid for network ports that are null or dot1q encapsulated. For QinQ network ports, the available labels are deducted by one.

⁸ Indicates the number of labels that the system assumes are always used on a specific service

⁹ Indicates the number of labels that the system subtracts from the total only if they are configured on the service

¹⁰ Number of labels that the system deducts from the total number

¹¹ Indicates a different number depending on the service type and the inner encapsulation used by each service, which reduces the maximum number of labels to push on egress

¹² Maximum SR-TE labels that the system can push when sending service packets on the wire

2.2.16.3 Flexible SR-TE label stack allocation for services

The 7705 SAR Gen 2 supports a dynamic egress label limit configuration mode that extends the number of allowed MPLS labels in the egress label stack by not counting specific labels in the BGP next-hop resolution check when those labels are not used. The configuration mode exists in EVPN services configured on Epipe, VPLS, and VPRN (EVPN-IFL), and in IP-VPN services.



Note: The 7705 SAR Gen 2 does not support ESI labels (EVPN multi-homing). ESI label information is included in this section for reference only.

Use the following commands to enable the dynamic egress label limit:

```
configure service vprn bgp-ipvprn mpls dynamic-egress-label-limit
configure service epipe bgp-evpn mpls dynamic-egress-label-limit
configure service vprn bgp-evpn mpls dynamic-egress-label-limit
configure service vpls bgp-evpn mpls dynamic-egress-label-limit
```

When the **dynamic-egress-label-limit** command is configured, the always-computed labels are no longer considered when resolving the next hop of the route. As a result, the following rules apply to the specified services:

- For VPRN services, the OAM label is never computed. This is true whether the BGP next hop is being resolved over an auto-bind tunnel or an SDP in the **vprn>spoke-sdp** context. The dynamic mode is supported for EVPN-IFL and IP-VPN families.
- For EVPN (Epipe or VPLS) services with **dynamic-egress-label-limit** configured, the CW and ESI label are only computed if they are used.
 - In the case of the CW, the system reduces the egress label limit by one label when the CW is configured in the service. The CW is always accounted when the **dynamic-egress-label-limit** command is not configured.
 - When the **dynamic-egress-label-limit** command is configured, the ESI label is not accounted for in Epipes or VPLS services without an ES; however, the ESI label is always accounted if **dynamic-egress-label-limit** is not configured.

When **no dynamic-egress-label-limit** is configured, the behavior follows the procedures described in [Control plane handling of egress label stack limitations](#).

In summary, when the **dynamic-egress-label-limit** is configured, the total amount of labels (X) configured in $X = (\text{max-sr-labels } Y + \text{additional-frr-labels } Z)$ can go higher for EVPN and IP-VPN services.

The following table summarizes the required behavior.

Table 11: Egress label stack limits for BGP services based on dynamic-egress-label-limit

Features that reduce the Label Stack		dynamic-egress-label-limit disabled			dynamic-egress-label-limit enabled		
		IP-VPN (VPRN)	EVPN-IFL (VPRN)	EVPN VPLS EVPN Epipe	IP-VPN (VPRN)	EVPN-IFL (VPRN)	EVPN VPLS EVPN Epipe
Always computed	Service label	1	1	1	1	1	1

Features that reduce the Label Stack		dynamic-egress-label-limit disabled			dynamic-egress-label-limit enabled		
		IP-VPN (VPRN)	EVPN-IFL (VPRN)	EVPN VPLS EVPN Epipe	IP-VPN (VPRN)	EVPN-IFL (VPRN)	EVPN VPLS EVPN Epipe
	OAM label	1	1	0	0	0	0
	CW	0	0	1	0	0	0
Computed if configured	Hash label (mutually exclusive with EL)	1	0	0	1	0	0
	EL+ELI	2	2	2	2	2	2
	CW	0	0	0	0	0	1
Required labels		2	2	2	1	1	1
Required labels + All Options		4	4	4	3	3	4
Maximum available labels		12	12	10	12	12	10
Maximum available transport labels without options		10	10	8	11	11	9
Maximum available transport labels with options		8	8	6	9	9	6

R-VPLS services, with EVPN-MPLS enabled, also support the **dynamic-egress-label-limit** command. When **dynamic-egress-label-limit** is configured, the CW is accounted for only if the **control-word** command is added.

- In R-VPLS services, the ESI label is not accounted because the routed encapsulation is always larger (and either ESI label for bridged traffic, or routed traffic without ESI label is transmitted by the R-VPLS).

2.2.17 IPv6 traffic engineering

This feature extends the traffic engineering capability with the support of IPv6 TE links and nodes.

This feature enhances IS-IS, BGP-LS, and the TE database with the additional IPv6 link TLVs and TE link TLVs. The following modes of operation are provided for IPv4 and IPv6 traffic engineering in a network.

- **Legacy mode**

This mode of operation enables the existing traffic engineering behavior for IPv4 RSVP-TE and IPv4 SR-TE. Only the RSVP-TE attributes are advertised in the legacy TE TLVs that are used by both RSVP-

TE and SR-TE LSP path computation in the TE domain routers. In addition, IPv6 SR-TE LSP path computation can now use these common attributes.

- **Legacy mode with application indication**

This mode of operation is intended for cases where link TE attributes are common to RSVP-TE and SR-TE applications and have the same value, but the user wants to indicate on a per-link basis which application is enabled.

Routers in the TE domain use these attributes to compute paths for IPv4 RSVP-TE LSP and IPv4/IPv6 SR-TE LSP.

- **Application-specific mode**

This mode of operation is intended for future use cases where TE attributes may have different values in RSVP-TE and SR-TE applications or are specific to one application (for example, RSVP-TE uses the Unreserved Bandwidth and Max Reservable Bandwidth attributes).

SR OS does not support configuring TE attributes that are specific to the SR-TE application. As a result, enabling this mode advertises the common TE attributes a single time using the Application Specific Link Attributes TLV. Routers in the TE domain use these attributes to compute paths for IPv4 RSVP-TE LSP and IPv4/IPv6 SR-TE LSP.

See [IS-IS IPv4/IPv6 SR-TE and IPv4 RSVP-TE feature behavior](#) for more information about the IPv4 and IPv6 Traffic Engineering modes of operation.

This feature also adds support of IPv6 destinations to the SR-TE LSP configuration. In addition, this feature extends the MPLS path configuration with hop indexes that include IPv6 addresses.

IPv6 SR-TE LSP is supported with the hop-to-label, local CSPF, and PCE path computation methods. This requires the enabling of the IPv6 traffic engineering feature in IS-IS.

2.2.17.1 Global configuration

To enable IPv6 TE on the router, the IPv6 TE router ID must have a valid IPv6 address. Use the following CLI command to configure the IPv6 TE router ID:

```
configure>router>ipv6-te-router-id interface interface-name
```

The IPv6 TE router ID is a mandatory parameter that uniquely identifies the router as being IPv6 TE capable to other routers in an IGP TE domain. IS-IS advertises this information using the IPv6 TE Router ID TLV as described in [TE attributes supported in IGP and BGP-LS](#).

When the command is not configured or the **no** form of the command is configured, the value of the IPv6 TE router ID reverts to the preferred primary global unicast address of the system interface. The user can also explicitly enter the name of the system interface to achieve the same outcome.

In addition, the user can specify a different interface and the preferred primary global unicast address of that interface is used instead. Only the system or a loopback interface is allowed because the TE router ID must use the address of a stable interface.

This address must be reachable from other routers in a TE domain and the associated interface must be added to IGP for reachability. Otherwise, IS-IS withdraws the advertisement of the IPv6 TE router ID TLV.

When configuring a new interface name for the IPv6 TE router ID, or when the same interface begins using a new preferred primary global unicast address, IS-IS immediately floods the new value.

If the referenced system is shut down or the referenced loopback interface is deleted or is shut down, or the last IPv6 address on the interface is removed, IS-IS withdraws the advertisement of the IPv6 TE router ID TLV.

2.2.17.2 IS-IS configuration

To enable the advertisement of additional link IPv6 and TE parameters, use the following **traffic-engineering-options** CLI syntax.

```
configure
router
  ipv6-te-router-id interface interface-name
  no ipv6-te-router-id
  [no] isis [instance]
    traffic-engineering
    no traffic-engineering
    traffic-engineering-options
    no traffic-engineering-options
    ipv6
    no ipv6
    application-link-attributes
    no application-link-attributes
    legacy
    no legacy
```

The **traffic-engineering** command is the main CLI command used to enable TE in an IS-IS instance. This command enables the advertisement of the IPv4 and TE link parameters using the legacy TE encoding, in accordance with RFC 5305. These parameters are used in IPv4 RSVP-TE and IPv4 SR-TE.

When the **ipv6** command under the **traffic-engineering-options** context is enabled, the traffic engineering behavior with IPv6 TE links is enabled. This IS-IS instance automatically advertises the new RFC 6119 IPv6 and TE TLVs and sub-TLVs, as described in [TE attributes supported in IGP and BGP-LS](#).

The **application-link-attributes** context allows the advertisement of the TE attributes of each link on a per-application basis. Two applications are supported in SR OS: RSVP-TE and SR-TE. The legacy mode of advertising TE attributes that is used in RSVP-TE is supported. Use the **no legacy** command to disable the legacy mode and also enable the per-application TE attribute advertisement for RSVP-TE.

See [IS-IS IPv4/IPv6 SR-TE and IPv4 RSVP-TE feature behavior](#) for more information about feature behavior and the interaction of the preceding CLI commands.

2.2.17.3 MPLS configuration

The SR-TE LSP configuration can accept an IPv6 address in the **to** and **from** parameters.

In addition, the MPLS path configuration can accept a hop index with an IPv6 address. The IPv6 address used in the **from** and **to** commands in the IPv6 SR-TE LSP, as well as the address used in the **hop** command of the path used with the IPv6 SR-TE LSP must correspond to the preferred primary global unicast IPv6 address of a network interface or a loopback interface of the corresponding LER or LSR router. The IPv6 address can also be set to the system interface IPv6 address. Failure to follow the preceding IPv6 address guidelines for the **from**, **to**, and **hop** commands causes path computation to fail with failure code "noCspfRouteToDestination".

A link-local IPv6 address of a network interface is also not allowed in the **hop** command of the path used with the IPv6 SR-TE LSP.

All other MPLS-level, LSP-level, and primary or secondary path-level configuration parameters available for an IPv4 SR-TE LSP are supported.

2.2.17.4 IS-IS, BGP-LS, and TE database extensions

IS-IS control plane extensions add support for the following RFC 6119 TLVs in IS-IS advertisements and TE-DB.

- IPv6 interface Address TLV (ISIS_TLV_IPv6_IFACE_ADDR 0xe8)
- IPv6 Neighbor Address sub-TLV (ISIS_SUB_TLV_NBR_IPADDR6 0x0d)
- IPv6 Global Interface Address TLV (only used by ISIS in IIH PDU)
- IPv6 TE Router ID TLV
- IPv6 SRLG TLV

IS-IS also supports the use of the Application Specific Link Attributes (ASLA) sub-TLV to advertise the protocol enabled on a specific TE-link (SR-TE, RSVP-TE, or both), as defined in *draft-ietf-isis-te-app*. The advertising router sends potentially different link TE attributes for RSVP-TE and SR-TE applications, and the router receiving the link TE attributes can identify the application that is enabled on the advertising router. For backward compatibility, the router continues to support the legacy mode of advertising link TE attributes, as recommended in RFC 5305, but the user can disable this mode.



Note: SR OS does not support configuring and advertising different link TE attribute values for RSVP-TE and SR-TE applications. The router advertises the same link TE attributes for both RSVP-TE and SR-TE applications.

See [IS-IS IPv4/IPv6 SR-TE and IPv4 RSVP-TE feature behavior](#) for more information about the behavior of the per-application TE capability.

These TLVs and sub-TLVs are advertised in IS-IS and added into the local TE-DB when received from IS-IS neighbors. In addition, if the **database-export** command is enabled in this ISIS instance, this information is also added in the Enhanced TE-DB.

The following enhancements are added to support advertisement of TE parameters in BGP-LS routes over an IPv4 or IPv6 transport.

- Importing IPv6 TE link TLVs from a local Enhanced TE-DB into the local BGP process for exporting to other BGP peers using the BGP-LS route family that is enabled on an IPv4 or an IPv6 transport BGP session
 - RFC 6119 IPv6 and TE TLVs and sub-TLVs are carried in the BGP-LS link NLRI, as per RFC 7752.
 - When the link TE attributes are advertised by IS-IS on a per-application basis using the ASLA TLV (ISIS TLV Type 16), they are carried in the new BGP-LS ASLA TLV (TLV Type TBD), in accordance with *draft-ietf-idr-bgp-ls-app-specific-attr*.
 - When a TE attribute of a link is advertised for both RSVP-TE and SRTE applications, there are three methods IS-IS can use. Each method results in a specific way the BGP-LS originator carries this information. These methods are summarized here, but more information is provided in [IS-IS IPv4/IPv6 SR-TE and IPv4 RSVP-TE feature behavior](#).
 - In the legacy mode of operation, all TE attributes are carried in the legacy IS-IS TE TLVs and the corresponding BGP-LS link attributes TLVs, as listed in [Table 12: Legacy link TE TLV support in TE-DB and BGP-LS](#).
 - In the legacy with application indication mode of operation, IGP and BGP-LS advertises the legacy TE attribute TLVs, and also advertises the ASLA TLV with the Legacy (L) flag set and the RSVP-TE and SR-TE application flags set. No TE sub-sub-TLVs are advertised within the ASLA TLV.

The legacy with application indication mode is intended for cases where link TE attributes are common to RSVP-TE and SR-TE applications and have the same value, but the user wants to indicate on a per-link basis which application is enabled.

- In the application-specific mode of operation, the TE attribute TLVs are sent as sub-sub-TLVs within the ASLA TLV. Common attributes to RSVP-TE and SR-TE applications have the main TLV L-flag cleared and the RSVP-TE and SR-TE application flags set. Any attribute that is specific to an application (RSVP-TE or SR-TE) is advertised in a separate ASLA TLV with the main TLV L-flag cleared and the specific application (RSVP-TE or SR-TE) flags set.

The application-specific mode of operation is intended for future cases where TE attributes may have different values in RSVP-TE and SR-TE applications or are specific to one application (for example, RSVP-TE uses the Unreserved Bandwidth and Max Reservable Bandwidth attributes).

- Exporting of any IPv6 and TE link TLVs that have been received from a BGP peer from the local BGP process to the local Enhanced TE-DB via a BGP-LS route family that is enabled on an IPv4 or IPv6 transport BGP session
- Exporting of IPv6 and TE link TLVs from local Enhanced TE-DB to NSP via the CPROTO channel on the VSR-NRC

2.2.17.4.1 BGP-LS originator node handling of TE attributes

The specification of the BGP-LS originator node in support of the ASLA TLV addresses the following main objectives.

1. Accommodate IGP node advertising the TE attribute in both legacy or application specific modes of operation.
2. Allow BGP-LS consumers (for example, PCE) that support the ASLA TLV to receive per-application attributes, even if the attribute values are duplicate, and easily store them per-application in the TE-DB. Also, if BGP-LS consumers receive legacy attributes, they can make a determination without ambiguity that these attributes are only for the RSVP-TE LSP application.
3. Provide continued support for older BGP-LS consumers that rely only on the legacy attributes.

The preceding objectives are supported by enhancements implemented in SR OS on the BGP-LS originator node. The following excerpts adapted from *draft-ietf-idr-bgp-ls-app-specific-attr* describe the enhancements:

1. Application-specific link attributes received from an IGP node without the use of ASLA encodings continue to be encoded using the respective BGP-LS top-level TLVs.
2. Application-specific link attributes received from an OSPF node using ASLA sub-TLV or from an IS-IS node using either ASLA sub-TLV or Application-Specific SRLG TLV must be encoded in the BGP-LS ASLA TLV as sub-TLVs. Exceptions to this rule are specified in 3.f and 3.g.
3. In the case of IS-IS, the following specific procedures are to be followed:
 - a. When application-specific link attributes are received from a node with the L-flag set in the IS-IS ASLA sub-TLV and application bits other than RSVP-TE are set in the application bit masks, the application-specific link attributes advertised in the corresponding legacy IS-IS TLVs/sub-TLVs must be encoded within the BGP-LS ASLA TLV as sub-TLVs with the application bits, other than the RSVP-TE bit, copied from the IS-IS ASLA sub-TLV. The link attributes advertised in the legacy IS-IS TLVs/sub-TLVs are also advertised in BGP-LS top-level TLVs as per [RFC 7752] [RFC 8571] [RFC 9104]. The same procedure also applies for the advertisement of the SRLG values from the IS-IS Application-Specific SRLG TLV.

- b. When the IS-IS ASLA sub-TLV has the RSVP-TE application bit set, the link attributes for the corresponding IS-IS ASLA sub-TLVs must be encoded using the respective BGP-LS top-level TLVs as per [RFC 7752] [RFC 8571] [RFC 9104]. Similarly, when the IS-IS Application-Specific SRLG TLV has the RSVP-TE application bit set, the SRLG values within it must be encoded using the top-level BGP-LS SRLG TLV (1096) as per [RFC 7752].
- c. The SRLGs advertised in IS-IS Application-Specific SRLG TLVs and the other link attributes advertised in IS-IS ASLA sub-TLVs are required to be collated, on a per-application basis, only for those applications that meet all of the following criteria:
 - Their bit is set in the SABM/UDABM in one of the two types of IS-IS encodings (for example, IS-IS ASLA sub-TLV).
 - The other encoding type (for example, IS-IS Application Specific SRLG TLV) has an advertisement with zero-length application bit masks.
 - There is no corresponding advertisement of that other encoding type (following the example, IS-IS Application Specific SRLG TLV) with that specific application bit set.

For each such application, its collated information must be carried in a BGP-LS ASLA TLV with that application's bit set in the SABM/UDABM.

- d. If the resulting set of collated link attributes and SRLG values is common across multiple applications, they may be advertised in a common BGP-LS ASLA TLV instance, where the bits for all such applications would be set in the application bit mask.
- e. Both the SRLG values from IS-IS Application-Specific SRLG TLVs and the link attributes from IS-IS ASLA sub-TLVs, with the zero-length application bit mask, must be advertised into a BGP-LS ASLA TLV with a zero-length application bit mask, independent of the collation described in 3.c and 3.d.
- f. [RFC 8919] allows the advertisement of the Maximum Link Bandwidth within an IS-IS ASLA sub-TLV, even though it is not an application-specific attribute. However, when originating the Maximum Link Bandwidth into BGP-LS, the attribute must be encoded only in the top-level BGP-LS Maximum Link Bandwidth TLV (1089) and must not be advertised within the BGP-LS ASLA TLV.
- g. [RFC 8919] also allows the advertisement of the Maximum Reservable Link Bandwidth and the Unreserved Bandwidth within an IS-IS ASLA sub-TLV, even though these attributes are specific to RSVP-TE application. However, when originating the Maximum Reservable Link Bandwidth and Unreserved Bandwidth into BGP-LS, these attributes must be encoded only in the BGP-LS top-level Maximum Reservable Link Bandwidth TLV (1090) and Unreserved Bandwidth TLV (1091) respectively and not within the BGP-LS ASLA TLV.

2.2.17.4.2 TE attributes supported in IGP and BGP-LS

[Table 12: Legacy link TE TLV support in TE-DB and BGP-LS](#) lists the TE attributes that are advertised using the legacy link TE TLVs defined in RFC 5305 for IS-IS and in RFC 3630 for OSPF. These TE attributes are carried in BGP-LS in accordance with RFC 7752. These legacy TLVs are already supported in SR OS and in IS-IS, OSPF, and BGP-LS.

To support IPv6 TE, the IS-IS IPv6 TE attributes (IPv6 TE router ID and IPv6 SRLG TLV) are advertised in BGP-LS in accordance with RFC 7752. These attributes can now be advertised within the ASLA TLV in IS-IS as recommended in RFC 8919 and in BGP-LS as recommended in *draft-ietf-idr-bgp-ls-app-specific-attr*. In the latter case, BGP-LS uses the same TLV type defined in RFC 7752 but is included as a sub-TLV of the new BGP-LS ASLA TLV. The following table also lists the code points for IS-IS and BGP-LS TLVs.

Table 12: Legacy link TE TLV support in TE-DB and BGP-LS

Link TE TLV description	IS-IS TLV type (RFC 5305)	OSPF TLV type (RFC 3630)	BGP-LS link NLRI link-attribute TLV type (RFC 7752)
Administrative group (color)	3	9	1088
Maximum link bandwidth	9	6	1089
Maximum reservable link bandwidth	10	7	1090
Unreserved bandwidth	11	8	1091
TE Default Metric	18	5	1092
SRLG	138 (RFC 4205)	16 (RFC 4203)	1096
IPv6 SRLG TLV	139 (RFC 6119)	—	1096
IPv6 TE Router ID	140 (RFC 6119)	—	1029
Application Specific Link Attributes	16 (RFC 8919)	—	1122 (provisional, as per <i>draft-ietf-idr-bgp-ls-app-specific-attr</i>)
Application Specific SRLG TLV	238 (RFC 8919)	—	1122 (provisional, as per <i>draft-ietf-idr-bgp-ls-app-specific-attr</i>)

The following table lists the TE attributes that are received from a third-party router implementation in legacy TE TLVs, or in the ASLA TLV for the RSVP-TE or SR-TE applications that are added into the local SR OSTE-DB; these are also distributed by the BGP-LS originator. However, these TLVs are not originated by an SR OS router IGP implementation.

Table 13: Additional link TE TLV support in TE-DB and BGP-LS

Link TE TLV description	IS-IS TLV type (RFC 7810)	OSPF TLV type (RFC 7471)	BGP-LS link NLRI link-attribute TLV type (RFC 8571)
Unidirectional Link Delay	33	27	1114
Min/Max Unidirectional Link Delay	34	28	1115
Unidirectional Delay Variation	35	29	1116
Unidirectional Link Loss	36	30	1117
Unidirectional Residual Bandwidth	37	31	1118
Unidirectional Available Bandwidth	38	32	1119

Link TE TLV description	IS-IS TLV type (RFC 7810)	OSPF TLV type (RFC 7471)	BGP-LS link NLRI link-attribute TLV type (RFC 8571)
Unidirectional Utilized Bandwidth	39	33	1120

Any other TE attribute received in a legacy TE TLV or in an Application Specific Link Attributes TLV is not added to the local router TE-DB and, therefore, is not distributed by the BGP-LS originator.

2.2.17.4.3 Support of MT-ISIS MT2 TE link attributes and router capability

SR OS supports the advertisement of existing MT-ISIS MT0 TE attributes in MT-ISIS MT2 and MT-ISIS MT2 export with BGP-LS.

Advertising TE link attributes with MT-ISIS MT2 depends on enabling SR-MPLS TE for IPv6 in MT0; an IPv6 TE router ID is encoded, even though TE is supported in MT-ISIS MT0 only. Use the following command to enable MT2.

```
configure router isis segment-routing multi-topology mt2
```

When MT2 is enabled and when delay is advertised, sub-TLVs (legacy and Application Specific Link Attributes [ASLA]) are distributed into MT-ISIS MT2. Use the following command to enable the advertisement of link delay in the IGP LSDB within legacy TE attributes in IS-IS or within ASLA when ASLA is enabled for SR-TE or RSVP-TE applications.

```
configure router isis traffic-engineering-options advertise-delay
```

Special attention is taken when leaking the IS-IS router capability when both MT0 and MT2 are enabled. When leaking router capability between IS-IS levels, as defined in RFC 7981, a reachability check must be performed. The Area Border Router (ABR) performs the following reachability check with MT-ISIS MT2 enabled:

- leak the router capability TLVs with a valid IPv4 router ID via IPv4 MT0
- for router capabilities with no valid IPv4 router ID but a valid IPv6 router ID, perform a reachability check via MT0 and MT2
- when either an IS-IS IPv4 or IPv6 router ID is reachable, then redistribute router capability is redistributed

When traffic engineering behavior with IPv6 links is enabled along with MT-ISIS MT2, this enables the advertisement of:

- MT2 IPv6 TLVs
- SR-MPLS TLVs
- Link TE TLVs with MT2

Use the following command to enable the advertisement of IPv6 TE in the IS-IS instance.

```
configure router isis traffic-engineering-options ipv6
```

SR OS supports advertising the Link TE attributes in MT2 legacy and in ASLA sub-TLV for the SR policy (SR-TE) application, as defined in RFC 8919.

Table 14: Legacy Link TE TLV support in TE-DB and BGP-LS

Link TE sub-TLV description	IS-IS TLV type (RFC 5305)	BGP-LS Link NLRI Link-Attribute TLV type (RFC 7752)
Administrative Group (color)	3	1088
Extended Administrative Group (EAG) (color) ¹³	14	1173
Maximum Link Bandwidth	9	1089
TE Default Metric	18	1092
IPv6 SRLG	139 (RFC 6119)	1096
IPv6 TE Router ID	140 (RFC 6119)	1029
Application Specific Link Attributes	16 (RFC 8919)	1122 (RFC 9294)
Min/Max Unidirectional Link Delay	34	1115

In addition, the following table shows TLVs that can be added into the local SR OS TE-DB and distributed by the BGP-LS originator. The following TLVs are originated by third-party routers and are not originated by an SR OS router IGP implementation. These TLVs can include TE attributes received from a third-party router implementation in legacy TE TLVs or in the ASLA TLV for the RSVP-TE.

Table 15: Additional Link TE TLV support in TE-DB and BGP-LS

Link TE TLV description	ISIS TLV type (RFC 7810)	OSPF TLV type (RFC 7471)	BGP-LS Link NLRI Link-Attribute TLV type (<i>draft-ietf-idr-te-pm-bgp</i>)
Unidirectional Link Delay	33	27	1114
Unidirectional Delay Variation	35	29	1116
Unidirectional Link Loss	36	30	1117
Unidirectional Residual Bandwidth	37	31	1118
Unidirectional Available Bandwidth	38	32	1119
Unidirectional Utilized Bandwidth	39	33	1120

¹³ The local router only advertises and supports EAG with Flexible Algorithms.

2.2.17.5 IS-IS IPv4/IPv6 SR-TE and IPv4 RSVP-TE feature behavior

The TE feature in IS-IS allows the advertising router to indicate to other routers in the TE domain which applications the advertising router has enabled: RSVP-TE, SR-TE, or both. As a result, a receiving router can safely prune links that are not enabled in one of the applications from the topology when computing a CSPF path in that application.

TE behavior consists of the following steps.

1. A valid IPv6 address value must exist for the system or loopback interface assigned to the **ipv6-te-router-id** command. The IPv6 address value can be either the preferred primary global unicast address of the system interface (default value) or that of a loopback interface (user-configured value).

The IPv6 TE router ID is mandatory for enabling IPv6 TE and enabling the router to be uniquely identified by other routers in an IGP TE domain as being IPv6 TE capable. If a valid value does not exist, then the IPv6 and TE TLVs described in [IS-IS, BGP-LS, and TE database extensions](#) are not advertised.

2. The **traffic-engineering** command enables the existing traffic engineering behavior for IPv4 RSVP-TE and IPv4 SR-TE. Enable the **rsvp** context on the router and enable **rsvp** on the interfaces to have IS-IS begin advertising TE attributes in the legacy TLVs. By default, the **rsvp** context is enabled as soon as the **mpls** context is enabled on the interface. If the **ipv6** command is also enabled, then the RFC 6119 IPv6 and TE link TLVs described above are advertised such that a router receiving these advertisements can compute paths for IPv6 SR-TE LSP in addition to paths for IPv4 RSVP-TE LSP and IPv4 SR-TE LSP. The receiving node cannot determine if IPv4 RSVP-TE, IPv4 SR-TE, or IPv6 SR-TE applications are enabled on the other routers. Legacy TE routers must assume that RSVP-TE is enabled on those remote TE links it received advertisements for.
3. When the **ipv6** command is enabled, IS-IS automatically begins advertising the RFC 6119 TLVs and sub-TLVs: the IPv6 TE router ID TLV, the IPv6 interface Address sub-TLV and Neighbor Address sub-TLV, or the Link-Local Interface Identifiers sub-TLV if the interface has no global unicast IPv6 address. The TLVs and sub-TLVs are advertised regardless of whether TE attributes are added to the interface in the **mpls** context. The advertisement of these TLVs is only performed when the **ipv6** command is enabled and **ipv6-routing** is enabled in this IS-IS instance and **ipv6-te-router-id** has a valid IPv6 address.

A network IP interface is advertised with the Link-Local Interface Identifiers sub-TLV if the network IP interface meets the following conditions:

- network IP interface has a link-local IPv6 address and no global unicast IPv6 address on the interface **ipv6** context
- network IP interface has no IPv4 address (and may or may not have the **unnumbered** option enabled on the interface **ipv4** context)

4. The **application-link-attributes** command enables the ability to send the link TE attributes on a per-application basis and explicitly conveys that RSVP-TE or SR-TE is enabled on that link on the advertising router.

Three modes of operation are allowed by the **application-link-attributes** command.

- [Legacy mode](#)
- [Legacy mode with application indication](#)
- [Application-specific mode](#)

The following table summarizes the IS-IS link TE parameter advertisement details for the three modes of operation of the IS-IS advertisement.

Table 16: Details of link TE advertisement methods

IGP traffic engineering options		Link TE advertisement details		
		RSVP-TE (rsvp enabled on interface)	SR-TE (segment-routing enabled in IGP instance)	RSVP-TE and SR-TE (rsvp enabled on interface and segment-routing enabled in IGP instance)
Legacy mode: Use the no application-link-attributes command.		Legacy TE TLVs	—	Legacy TE TLVs
Legacy mode with application indication: Enable configure router isis traffic-engineering-options application-link-attributes legacy	rsvp disabled on router (rsvp operationally down on all interfaces)	—	Legacy TE TLVs ASLA TLV -Flags: {Legacy=0, SR-TE=1}; TE sub-sub-TLVs	Legacy TE TLVs ASLA TLV -Flags: {Legacy=1, RSVP-TE=0, SR-TE=1}
	rsvp enabled on router	Legacy TE TLVs ASLA TLV -Flags: {Legacy=1, RSVP-TE=1}	Legacy TE TLVs ASLA TLV -Flags: {Legacy=1, SR-TE=1}	Legacy TE TLVs ASLA TLV -Flags: {Legacy=1, RSVP-TE=1, SR-TE=1}
Application specific mode: Disable configure router isis traffic-engineering-options application-link-attributes legacy		ASLA TLV -Flags: {Legacy=0, RSVP-TE=1}; TE sub-sub-TLVs	ASLA TLV -Flags: {Legacy=0, SR-TE=1}; TE sub-sub-TLVs	ASLA TLV -Flags: {Legacy=0, RSVP-TE=1; SR-TE=1}; TE sub-sub-TLVs (common attributes) ASLA TLV -Flags: {Legacy=0, RSVP-TE=1}; TE sub-sub-TLVs (RSVP-TE specific attributes; for example, Unreserved BW and Resvble BW) ASLA TLV -Flags: {Legacy=0, SR-TE=1}; TE sub-sub-TLVs (SR-TE specific attributes; not supported in SR OS 19.10.R1)

Legacy mode

For legacy mode, use the **no application-link-attributes** command.

The **application-link-attributes** command is disabled by default and the **no** form matches the TE behavior described in list item 2. It enables the existing TE behavior for IPv4 RSVP-TE and IPv4 SR-TE. Only the RSVP-TE attributes are advertised in the legacy TE TLVs that are used by both RSVP-TE and SR-TE LSP CSPF in the TE domain routers. No separate SR-TE attributes are advertised.

If the **ipv6** command is also enabled, then the RFC 6119 IPv6 and TE link TLVs are advertised in the legacy TLVs. A router in the TE domain receiving these advertisements can compute paths for IPv6 SR-TE LSP.

If the user shuts down the **rsvp** context on the router or on a specific interface, the legacy TE attributes of all the MPLS interfaces or of that specific MPLS interface are not advertised. Routers can still compute SR-TE LSPs using those links, but LSP path TE constraints are not enforced because the links appear in the TE Database as if they did not have TE parameters.

[Table 12: Legacy link TE TLV support in TE-DB and BGP-LS](#) shows the encoding of the legacy TE TLVs in both IS-IS and BGP-LS.

Legacy mode with application indication

To use legacy mode with application indication, enable the **legacy** command in the **configure router isis traffic-engineering-options application-link-attributes** context.

The legacy with application indication mode is intended for cases where link TE attributes are common to RSVP-TE and SR-TE applications and have the same value, but the user wants to indicate on a per-link basis which application is enabled.

IS-IS continues to advertise the legacy TE attributes for both RSVP-TE and SR-TE applications and includes the Application-Specific Link Attributes TLV with the application flag set to RSVP-TE, SR-TE, or both, but without the sub-sub-TLVs. IS-IS also advertises the Application-Specific SRLG TLV with the application flag set to RSVP-TE, SR-TE, or both, but without the actual values of the SRLGs.

Routers in the TE domain use these attributes to compute CSPF for IPv4 RSVP-TE LSP and IPv4 SR-TE LSP.

If the **ipv6** command is also enabled, the RFC 6119 IPv6 and TE TLVs are advertised. A router in the TE domain that receives these advertisements can compute paths for IPv6 SR-TE LSP.



Note: The **segment-routing** command must be enabled in the IS-IS instance; otherwise, the flag for the SR-TE application cannot be set in the Application-Specific Link Attributes TLV or Application-Specific SRLG TLV.

To disable advertising of RSVP-TE attributes, shut down the **rsvp** context on the router. However, doing so reverts to advertising the link SR-TE attributes using the Application-Specific Link Attributes TLV and the TE sub-sub-TLVs as shown in [Table 16: Details of link TE advertisement methods](#). If legacy attributes were used, legacy routers incorrectly interpret that this router enabled RSVP and may signal RSVP-TE LSP paths using its links.

[Table 12: Legacy link TE TLV support in TE-DB and BGP-LS](#) lists the code points for IS-IS and BGP-LS legacy TLVs.

Example

The following excerpt from the Link State Database (LSDB) shows the advertisement of TE parameters for a link with both RSVP-TE and SR-TE applications enabled.

```
TE IS Nbrs      :
  Nbr           : Dut-A.00
  Default Metric : 10
```

```

Sub TLV Len      : 124
IF Addr   : 10.10.2.3
IPv6 Addr : 3ffe::10:10:2:3
Nbr IP    : 10.10.2.1
Nbr IPv6  : 3ffe::10:10:2:1
MaxLink BW: 100000 kbps
Resvble BW: 500000 kbps
Unresvd BW:
    BW[0] : 500000 kbps
    BW[1] : 500000 kbps
    BW[2] : 500000 kbps
    BW[3] : 500000 kbps
    BW[4] : 500000 kbps
    BW[5] : 500000 kbps
    BW[6] : 500000 kbps
    BW[7] : 500000 kbps
Admin Grp : 0x1
TE Metric : 123
TE APP LINK ATTR :
    SABML-flags:Legacy SABM-flags:RSVP-TE SR-TE
Adj-SID: Flags:v4VL Weight:0 Label:524287
Adj-SID: Flags:v6BVL Weight:0 Label:524284
TE SRLGs :
    SRLGs : Dut-A.00
    Lcl Addr : 10.10.2.3
    Rem Addr : 10.10.2.1
    Num SRLGs : 1
    1003
TE APP SRLGs :
    Nbr : Dut-A.00
    SABML-flags:Legacy SABM-flags: SR-TE
    IF Addr : 10.10.2.3
    Nbr IP : 10.10.2.1

```

Application-specific mode

To use legacy mode with application indication, disable the **legacy** command in the **configure router isis traffic-engineering-options application-link-attributes** context.

The application-specific mode of operation is intended for future use cases where TE attributes may have different values in RSVP-TE and SR-TE applications (this capability is not supported in SR OS) or are specific to one application (for example, RSVP-TE uses the Unreserved Bandwidth and Max Reservable Bandwidth attributes).

IS-IS advertises the TE attributes that are common to RSVP-TE and SR-TE applications in the sub-sub-TLVs of the new ASLA sub-TLV. IS-IS also advertises the link SRLG values in the Application-Specific SRLG TLV. In both cases, the application flags for RSVP-TE and SR-TE are also set in the sub-TLV.

IS-IS advertises the TE attributes that are specific to the RSVP-TE application separately in the sub-sub-TLVs of the new application attribute sub-TLV. The application flag for RSVP-TE is also set in the sub-TLV.

SR OS does not support configuring and advertising TE attributes that are specific to the SR-TE application.

Common value RSVP-TE and SR-TE TE attributes are combined in the same application attribute sub-TLV with both application flags set, while the non-common value TE attributes are sent in their own application attribute sub-TLV with the corresponding application flag set.

[Figure 23: Attribute mapping per application](#) shows an excerpt from the Link State Database (LSDB). Attributes in green font are common to both RSVP-TE and SR-TE applications and are combined, while the attribute in red font is specific to the RSVP-TE application and is sent separately.

Figure 23: Attribute mapping per application

```

TE IS Nbrs      :
  Nbr      : Dut-A.00
  Default Metric : 100
  Sub TLV Len   : 111
  IF Addr      : 1.0.13.3
  IPv6 Addr     : 3ffe::102:606
  Nbr IP       : 1.0.13.1
  Adj-SID: Flags:v4BVL Weight:0 Label:524285
  Adj-SID: Flags:v6BVL Weight:0 Label:524284
  SABML-flags:Non-Legacy SABM-flags:RSVP-TE SR-TE
  MaxLink BW: 99999997 kbps
  Admin Grp   : 0x0
  TE Metric   : 100
  SABML-flags:Non-Legacy SABM-flags:RSVP-TE
  Resvble BW: 99999997 kbps
  Unresvd BW:
    BW[0] : 99999997 kbps
    BW[1] : 99999997 kbps
    BW[2] : 99999997 kbps
    BW[3] : 99999997 kbps
    BW[4] : 99999997 kbps
    BW[5] : 99999997 kbps
    BW[6] : 99999997 kbps
    BW[7] : 99999997 kbps
TE APP SRLGs    :
  Nbr : Dut-A.00
  SABML-flags:Non-Legacy SABM-flags:RSVP-TE SR-TE
  IF Addr : 1.0.13.3
  Nbr IP   : 1.0.13.1
  Num SRLGs : 1
  SRLGs    : 1

```

sw0973

Routers in the TE domain use these attributes to compute CSPF for IPv4 RSVP-TE LSPs and IPv4 SR-TE LSPs. If the **ipv6** command is also enabled, the RFC 6119 IPv6 TLVs are advertised. A router in the TE domain receiving these advertisements can compute paths for IPv6 SR-TE LSP.



Note: The **segment-routing** command must be enabled in the IS-IS instance or the common TE attribute is not advertised for the SR-TE application.

To disable advertising of RSVP-TE attributes, use the **rsvp shutdown** command on the router.

2.2.17.6 IPv6 SR-TE LSP support in MPLS

This feature is supported with the hop-to-label, the local CSPF, and the PCE (PCC-initiated and PCE-initiated) path computation methods.

IPv6 SR-TE LSP is supported in IS-IS MT0 or MT2 topology. The user must enable IPv6 forwarding in IS-IS MT0 or MT2 and enable the advertisement of IPv6, TE, and SR-MPLS link information in native or MT IS-IS TLVs. Without these, the local CSPF cannot compute a path, select links, and assign SIDs to the SR-TE LSP path to an IPv6 destination.

Use the following command to enable the IS-IS instance to advertise IPv6 link and prefix TLVs in MT0 (**native** value) or MT2 (**mt** value).

```
configure router isis ipv6-routing {native | mt}
```

Use the following command to enable unicast IPv6 forwarding in MT2 and advertise IPv6 link and prefix information in IS-IS MT TLVs.

```
configure router isis multi-topology ipv6-unicast
```

Use the following command to configure the IPv6 TE router ID. When advertised in IS-IS as per the next command (**traffic-engineering-options ipv6**), it identifies the router as IPv6-TE capable.

```
configure router ipv6-te-router-id interface interface-name
```

Use the following command to advertise the IPv6 TE router ID and the RFC 6119 IPv6 and TE link TLVs.

```
configure router isis traffic-engineering-options ipv6
```

Use the following command to enable MPLS forwarding in MT0 or MT2 (optional parameter **multi-topology mt2** must be configured]), and advertise prefix and link SR-MPLS information in IS-IS MT TLVs.

```
configure router isis segment-routing  
configure router isis segment-routing multi-topology mt2
```

All capabilities of an IPv4 provisioned SR-TE LSP are supported with an IPv6 SR-TE LSP unless otherwise indicated. For more information, see [SR-TE LSP path computation using hop-to-label translation](#) and [SR-TE LSP path computation using local CSPF](#).

The following describes some important differences between an IPv4 and IPv6 SR-TE LSP support in MPLS.

The IPv6 address used in the **from** and **to** commands in the IPv6 SR-TE LSP, as well as the address used in the **hop** command of the path used with the IPv6 SR-TE LSP, must correspond to the preferred primary global unicast IPv6 address of a network interface or a loopback interface of the corresponding LER or LSR router. The IPv6 address can also be set to the system interface IPv6 address. Failure to follow the preceding IPv6 address guidelines for the **from**, **to**, and **hop** commands causes path computation to fail with failure code "noCspfRouteToDestination". A link-local IPv6 address of a network interface is also not allowed in the **hop** command of the path used with the IPv6 SR-TE LSP. The configuration fails.

A TE link with no global unicast IPv6 address and only a link local IPv6 address can be used in the path computation by the local CSPF. The address shown in the Computed Hops and in the Actual Hops fields of the output of the path **show** command uses the neighbor's IPv6 TE router ID and the Link-Local Interface Identifiers sub-TLV. The exceptions are if the interface is of type broadcast or type point-to-point but also has a local IPv4 address. Only the neighbor's IPv6 TE router ID is shown, as the Link-Local Interface Identifiers sub-TLV is not advertised in these situations.

The UP value of the global MPLS IPv4 state requires that the system interface be in the admin UP state and to have a valid IPv4 address.

The UP value of the global MPLS IPv6 state requires that the interface used for the IPv6 TE router ID be in admin UP state and to have a valid preferred primary IPv6 global unicast address.

The UP value of the TE interface MPLS IPv4 state requires the interface be in the admin UP state in the **router** context and the global MPLS IPv4 state be in UP state.

The UP value of the TE interface MPLS IPv6 state requires the interface be in the admin UP state in the **router** context and the global MPLS IPv6 state be in UP state.



Note: IPv6 forwarding and IPv6 SR-MPLS forwarding should only be enabled in either MT0 or MT2. The following limitations apply when IPv6 forwarding and IPv6 SR-MPLS forwarding are enabled in both IS-IS MT0 and MT2:

- If the destination prefix of the IPv6 SR-TE LSP or the prefix of a loose or strict hop of the path definition of the LSP is reachable in both MT0 and MT2, CSPF computes candidate paths in both MT0 and MT2 and selects one path within each topology using the procedure described in [SR-TE LSP path computation using local CSPF](#). In the final selection step, the lowest metric path dictates whether the MT0 or MT2 path is selected. If the metric is equal, the MT0 path is selected.
- IS-IS assigns both a protected SID and an unprotected SID to each MT2 adjacency between neighbors. However, IS-IS only assigns an unprotected SID for the MT0 adjacency. As a result, an IPv6 SR-TE LSP with the **configure router mpls lsp local-sr-protection** command set to **none** (use unprotected adjacencies only) will fail to find a path in MT0 topology.
- If the user changes the configuration of any of the following commands, IS-IS restarts all interfaces, which causes the IPv6 and IPv4 adjacency SIDs of each interface to be deprogrammed and reprogrammed in the data path. This causes the IPv4 and IPv6 SR-ISIS tunnels, as well as the SR-TE LSPs using these adjacencies, to flap.

```
configure router isis ipv6-routing {native | mt}
configure router isis multi-topology ipv6-unicast
configure router isis segment-routing multi-topology mt2
```

2.2.17.6.1 IPv6 SR-TE auto-LSP

This feature supports the auto-creation of an IPv6 SR-TE mesh LSP and for an IPv6 SR-TE one-hop LSP.

The SR-TE mesh LSP feature specifically binds an LSP template of type **mesh-p2p-srte** with one or more IPv6 prefix lists. When the TE-DB discovers a router that has an IPv6 TE router ID matching an entry in the prefix list, it triggers MPLS to instantiate an SR-TE LSP to that router using the LSP parameters in the LSP template.

The SR-TE one-hop LSP feature specifically activates an LSP template of type **one-hop-p2p-srte**. In this case, the TE database keeps track of each TE link that comes up to a directly connected IGP TE neighbor. It then instructs MPLS to instantiate an SR-TE LSP with the following parameters:

- the source IPv6 address of the local router
- an outgoing interface matching the interface index of the TE-link
- a destination address matching the IPv6 TE router ID of the neighbor on the TE link

A **family** CLI leaf is added to the LSP template configuration and must be set to the **ipv6** value. By default, this command is set to the **ipv4** value for backward compatibility. When establishing both IPv4 and IPv6 SR-TE mesh auto-LSPs with the same parameters and constraints, a separate LSP template of type **mesh-p2p-srte** must be configured for each address family with the **family** command set to the IPv4 or IPv6 value. SR-TE one-hop auto-LSPs can only be established for either IPv4 or IPv6 family, not both. The **family** command in the LSP template of type **one-hop-p2p-srte** should be set to the needed IP family value.



Note: An IPv6 SR-TE auto-LSP can be reported to a PCE but cannot be delegated or have its paths computed by the PCE.

All capabilities of an IPv4 SR-TE auto-LSP are supported with an IPv6 SR-TE auto-LSP unless indicated otherwise.

2.2.18 OSPF link TE attribute reuse

This section describes the support of OSPF application-specific TE link attributes.

2.2.18.1 OSPF application-specific TE link attributes

Existing definitions for the advertisement of OSPFv2 TE-related link attributes (for example, bandwidth) are used in RSVP-TE deployments (see *draft-ietf-spring-segment-routing-policy-07.txt* for more information). Initially, these TE-related link attributes were only used by RSVP-TE. However, additional applications emerged that also required link attributes (for example, SR-TE). The link attributes used by these new applications are not always identical to those advertised in RSVP-TE.

The usage of link attributes has introduced ambiguity in deployments that include a mix of RSVP-TE and SR-TE support. For example, it is not possible to unambiguously indicate the specific advertisements used by RSVP-TE and SR-TE. Although this may not be an issue for fully congruent topologies, any incongruence causes ambiguity. An additional issue arises in cases where both applications are supported on a link but the link attribute values associated with each application differ. Advertisements without OSPFv2 application-specific TE link attributes do not support the advertisement of application specific values for the same attribute on a specific link.

CLI syntax:

```
Config
router
  ospf
    traffic-engineering-options
      sr-te {legacy | application-specific-link-attributes}
    no sr-te
```

The **traffic-engineering-options** command enables the context to configure advertisement of the TE attributes of each link on a per-application basis. Two applications are supported in SR OS: RSVP-TE and SR-TE.

The **legacy** mode of advertising TE attributes that is used in RSVP-TE is still supported. In addition, the following configuration options are allowed:

- **no sr-te**

This option advertises the TE information for RSVP links using TE opaque LSAs. The **no** form is the default value.

- **sr-te legacy**

This option advertises the TE information for MPLS-enabled SR links using TE opaque LSAs.



Note: The operator should not use the **sr-te legacy** option if the network has both RSVP-TE and SR-TE applied and the links are not congruent.

- **sr-te application-specific-link-attributes**

This option advertises the TE information for MPLS-enabled SR links using the new Application Specific Link Attributes (ASLA) TLVs.

See RFC 8920 for definitions of a subset of all possible TE extensions and TE Metric Extensions that can be encoded within Application Specific Link sub-TLVs. The following table describes the relevant values for SR OS.

Table 17: Nokia support for ASLA extended link TLV encoding

OSPFv2 extended link TLV sub-TLVs (RFC 7684)				
IANA	Attribute type	TE-DB ¹⁴	SR OS sub-TLV of Extended Link TLV ¹⁵	SR OS Nested sub-TLV of ASLA Extended Link TLV encoding ¹⁶
10	ASLA	✓	✓	—
11	Shared Risk Link Group	✓	—	✓
12	Unidirectional Link Delay	✓	—	—
13	Min/Max Unidirectional Link Delay	✓	—	—
14	Unidirectional Delay Variation	✓	—	—
15	Unidirectional Link Loss	✓	—	—
16	Unidirectional Residual Bandwidth	✓	—	—
17	Unidirectional Available Bandwidth	✓	—	—
18	Unidirectional Utilized Bandwidth	✓	—	—
19	Administrative Group	✓	—	✓
20	Extended Administrative Group	✓	—	—
22	TE Metric	✓	—	✓
23	Maximum Link Bandwidth	✓	✓	—

¹⁴ Support to include the attributes from received LSAs into the Nokia TE-DB and export into BGP-LS. See *draft-ietf-idr-bgp-ls-app-specific-attr* for more information.

¹⁵ Node support to encode the link attribute as a sub-TLV in an OSPFv2 Extended Link TLV.

¹⁶ Node support to encode the link attribute as a sub-TLV in an OSPFv2 Application Specific Extended Link sub-TLV.

The solution proposed in *draft-ietf-ospf-te-link-attr-reuse-14* assumes that OSPF does not need to move all RSVP-TE attributes from the TE Opaque LSA into the extended link LSA. For, RSVP-TE, consequently, there is no significant modification and it can continue to be advertised using existing OSPF TLVs. For SR-TE and future applications, the ASLA TLVs may be used. Alternatively, existing TE Opaque LSAs could be used through configuration. The following table describes the possible configurations for TE Opaque LSAs.

Table 18: Configuration considerations for TE Opaque LSAs

IGP configuration	ospf>traffic-engineering <20.7	ospf>traffic-engineering ospf>te-opts>no sr-te	ospf>traffic-engineering ospf>te-opts>sr-te legacy	ospf>traffic-engineering ospf>te-opts>sr-te application-link-attribute
Interface configuration	—	—	—	—
MPLS + RSVP	TE-Opaque	TE-Opaque	TE-Opaque	TE-Opaque
MPLS + SR	—	—	TE-Opaque ¹⁷	ASLA (SR-TE)
MPLS + RSVP + SR	TE-Opaque	TE-Opaque	TE-Opaque	TE-Opaque (RSVP) + ASLA (SR-TE)

2.2.19 Configuring and operating SR-TE

This section provides information about the configuration and operation of the SR-TE LSP.

2.2.19.1 SR-TE configuration prerequisites

About this task

To configure SR-TE, the user must first configure prerequisite parameters.

Procedure

- Step 1.** Configure the label space partition for the Segment Routing Global Block (SRGB) for all participating routers in the segment routing domain by using the **mpls-labels>sr-labels** command.

Example

```
mpls-labels
- sr-labels start 200000 end 200400
- exit
```

- Step 2.** Enable segment routing, traffic engineering, and advertisement of router capability in all participating IGP instances in all participating routers by using the **traffic-engineering**, **advertise-router-capability**, and **segment-routing** commands.

¹⁷ If the local router interface is configured with MPLS and SR, and RSVP-TE is deployed on remote servers, the remote routers incorrectly conclude that the link is RSVP-enabled.

Example

```
ospf 0
- traffic-engineering
- advertise-router-capability area
- loopfree-alternates remote-lfa
- area 0.0.0.202
  - stub
    - no summaries
  - exit
- interface "system"
  - node-sid index 194
  - no shutdown
- exit
- interface "toSim199"
  - interface-type point-to-point
  - no shutdown
- exit
- interface "toSim213"
  - interface-type point-to-point
  - no shutdown
- exit
- interface "toSim219"
  - interface-type point-to-point
  - metric 2000
  - no shutdown
- exit
- exit
- segment-routing
  - prefix-sid-range global
  - no shutdown
- exit
- no shutdown
- exit
```

Step 3. Configure an segment routing tunnel MTU for the IGP instance, if required, by using the **tunnel-mtu** command.

Example

```
prefix-sid-range global
- tunnel-mtu 1500
- no shutdown
```

Step 4. Assign a node SID to each loopback interface that a router would use as the destination of a segment routing tunnel by using the **node-sid** command.

Example

```
ospf 0
- area 0.0.0.202
  - interface "system"
    - node-sid index 194
    - no shutdown
  - exit
```

2.2.19.2 SR-TE LSP configuration overview

The user can configure an SR-TE LSP as a label switched path (LSP) under the MPLS context by specifying the **sr-te** LSP type.

```
config>router>mpls>lsp lsp-name [mpls-tp src-tunnel-num | sr-te]
```

The user can configure a primary path for an RSVP LSP.

Use the following CLI syntax to associate an empty path or a path with strict or loose explicit hops with the primary paths of the SR-TE LSP:

```
config>router>mpls>path>hop hop-index ip-address {strict | loose}
- config>router>mpls>lsp>primary path-name
```

2.2.19.3 Configuring SR-TE LSP label stack size

Use the following command to configure the maximum number of labels that the ingress LER can push for a specific SR-TE LSP.

```
configure router mpls lsp max-sr-labels
```

This command allows the user to reduce the SR-TE LSP label stack size by accounting for additional transport, service, and other labels when packets are forwarded in a particular context. See [Datapath support](#) for more information about label stack size requirements in various forwarding contexts. If the CSPF on the PCE or the hop-to-label translation of the router cannot find a path that meets the maximum SR label stack, the SR-TE LSP remains on its current path or remains down if it has no path. The range is 1-11 labels with a default value of 6.

2.2.19.4 Configuring adjacency SID parameters

Configure the adjacency hold timer for the LFA or remote LFA backup next hop of an adjacency SID.

Use the following CLI command syntax to configure the length of the interval during which LTN or ILM records of an adjacency SID are kept:

```
config>router>ospf>segment-routing>adj-sid-hold seconds[1..300, default 15]
- config>router>isis>segment-routing>adj-sid-hold seconds[1..300, default 15]
```

```
adj-sid-hold 15
- no entropy-label-capability
- prefix-sid-range global
- no tunnel-table-pref
- no tunnel-mtu
- no backup-node-sid
- no shutdown
```

When protection is enabled globally for all node SIDs and local adjacency SIDs with the **loopfree-alternates** command in IS-IS or OSPF at the LER and LSR, applications may exist for which the user wants traffic to never divert from the strict hop computed by CSPF for an SR-TE LSP. In such cases, use

the following CLI command syntax to disable protection for all adjacency SIDs formed over a network IP interface:

```
config>router>ospf>area>if>no sid-protection
- config>router>isis>if>no sid-protection
```

Example: Configuration output

```
node-sid index 194
- no sid-protection
- no shutdown
```

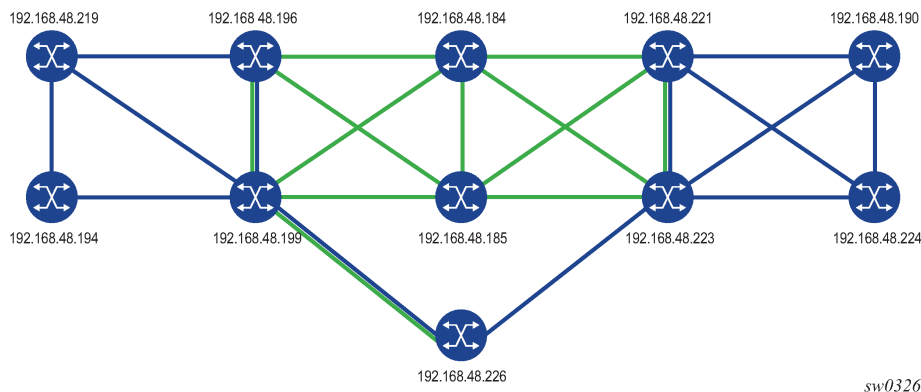
2.2.19.5 Configuring a mesh of SR-TE auto-LSPs

The following shows the detailed configuration for the creation of a mesh of SR-TE auto-LSPs. The network uses IS-IS with the backbone area being in Level 2 and the leaf areas being in Level 1.

The NSP is used for network discovery only and the NRC-P learns the network topology using BGP-LS.

The following figure shows the view of the multilevel IS-IS topology in the NSP GUI. The backbone Level 2 area is highlighted in green.

Figure 24: Multilevel IS-IS topology in the NSP GUI



The mesh of SR-TE auto-LSPs is created in the backbone area and originates on an ABR node with address 192.168.48.199 (Phoenix 199). The LSP template uses a default path that includes an anycast SID prefix corresponding to transit routers 192.168.48.184 (Dallas 184) and 192.168.48.185 (Houston 185).

Output example

The following is the configuration of transit router Dallas 184, which shows the creation of a loopback interface with the anycast prefix and the assignment of a SID to it. The same configuration must be performed on the transit router Houston 185. See lines marked with an asterisk (*).

```
*A:Dallas 184>config>router# info
-----
echo "IP Configuration"
#-----
    if-attribute
        admin-group "olive" value 20
```

```

        admin-group "top" value 10
        srlg-group "top" value 10
    exit
    interface "anycast-sid"
        address 192.168.48.99/32
        loopback
        no shutdown
    exit
    interface "system"
        address 192.168.48.184/32
        no shutdown
    exit
    interface "toJun164"
        address 10.19.2.184/24
        port 1/1/4:10
        no shutdown
    exit
    interface "toSim185"
        address 10.0.3.184/24
        port 1/1/2
        no shutdown
    exit
    interface "toSim198"
        address 10.0.2.184/24
        port 1/1/3
        if-attribute
            admin-group "olive"
        exit
        no shutdown
    exit
    interface "toSim199"
        address 10.0.13.184/24
        port 1/1/5
        no shutdown
    exit
    interface "toSim221"
        address 10.0.4.184/24
        port 1/1/1
        no shutdown
    exit
    interface "toSim223"
        address 10.0.14.184/24
        port 1/1/6
        no shutdown
    exit
#-----

*A:Dallas 184>config>router>isis# info
-----
    level-capability level-2
    area-id 49.0000
    database-export identifier 10 bgp-ls-identifier 10
    traffic-engineering
    advertise-router-capability area
    level 2
        wide-metrics-only
    exit
    interface "system"
        ipv4-node-sid index 384
        no shutdown
    exit
    interface "toSim198"
        interface-type point-to-point
        no shutdown

```

```

exit
interface "toSim185"
    interface-type point-to-point
    no shutdown
exit
interface "toSim221"
    interface-type point-to-point
    no shutdown
exit
interface "toSim199"
    interface-type point-to-point
    level 2
    metric 100
exit
no shutdown
exit
interface "toSim223"
    interface-type point-to-point
    level 2
    metric 100
exit
no shutdown
exit
interface "anycast-sid"
    ipv4-node-sid index 99
    no shutdown
exit
segment-routing
    prefix-sid-range global
    no shutdown
exit
no shutdown
-----

```

In the ingress LER Phoenix 199 router, the anycast SID is learned from both transit routers, but is currently resolved in IS-IS to transit router Houston 185. See lines marked with an asterisk (*).

```

*A:Phoenix 199# show router isis prefix-sids
=====
Rtr Base ISIS Instance 0 Prefix/SID Table
=====

```

Prefix	SID	Lvl/Typ	SRMS MT	AdvRtr Flags	
192.168.48.194/32	399	1/Int.	N	Reno 194	
			0	NnP	
192.168.48.194/32	399	2/Int.	N	Salt Lake 198	
			0	RNnP	
192.168.48.194/32	399	2/Int.	N	Phoenix 199	
			0	RNnP	
192.168.48.99/32	99	2/Int.	N	Dallas 184	*
			0	NnP	*
192.168.48.99/32	99	2/Int.	N	Houston 185	*
			0	NnP	*
192.168.48.184/32	384	2/Int.	N	Dallas 184	
			0	NnP	
192.168.48.185/32	385	2/Int.	N	Houston 185	
			0	NnP	
192.168.48.190/32	390	2/Int.	N	Chicago 221	
			0	RNnP	
192.168.48.190/32	390	2/Int.	N	St Louis 223	
			0	RNnP	
192.168.48.194/32	394	1/Int.	N	Reno 194	

```

192.168.48.194/32      394      2/Int.      N      0      NnP      Salt Lake 198
192.168.48.194/32      394      2/Int.      N      0      RnNP
192.168.48.194/32      394      2/Int.      N      0      Phoenix 199
192.168.48.198/32      398      1/Int.      N      0      RnNP      Salt Lake 198
192.168.48.198/32      398      2/Int.      N      0      NnP      Salt Lake 198
192.168.48.198/32      398      2/Int.      N      0      NnP      Phoenix 199
192.168.48.199/32      399      2/Int.      N      0      RnNP      Salt Lake 198
192.168.48.199/32      399      1/Int.      N      0      RnNP      Phoenix 199
192.168.48.199/32      399      2/Int.      N      0      NnP      Phoenix 199
192.168.48.219/32      319      2/Int.      N      0      NnP      Salt Lake 198
192.168.48.219/32      319      2/Int.      N      0      RnNP      Phoenix 199
192.168.48.219/32      319      1/Int.      N      0      RnNP      Las Vegas 219
192.168.48.221/32      321      2/Int.      N      0      NnP      Chicago 221
192.168.48.221/32      321      2/Int.      N      0      NnP      St Louis 223
192.168.48.223/32      323      2/Int.      N      0      RnNP      Chicago 221
192.168.48.223/32      323      2/Int.      N      0      RnNP      St Louis 223
192.168.48.224/32      324      2/Int.      N      0      NnP      Chicago 221
192.168.48.224/32      324      2/Int.      N      0      RnNP      St Louis 223
192.168.48.226/32      326      2/Int.      N      0      RnNP      PCE Server 226
3ffe::a14:194/128      294      1/Int.      N      0      NnP      Reno 194
3ffe::a14:194/128      294      2/Int.      N      0      NnP      Phoenix 199
3ffe::a14:199/128      299      1/Int.      N      0      RnNP      Phoenix 199
3ffe::a14:199/128      299      2/Int.      N      0      NnP      Phoenix 199
-----
No. of Prefix/SIDs: 32 (15 unique)
-----
SRMS : Y/N = prefix SID advertised by SR Mapping Server (Y) or not (N)
      S    = SRMS prefix SID is selected to be programmed
Flags: R    = Re-advertisement
      N    = Node-SID
      nP   = no penultimate hop POP
      E    = Explicit-Null
      V    = Prefix-SID carries a value
      L    = value/index has local significance
=====
*A:Phoenix 199# tools dump router segment-routing tunnel
=====
Legend: (B) - Backup Next-hop for Fast Re-Route
        (D) - Duplicate
=====
-----+

```

Prefix Sid-Type	Fwd-Type Next Hop(s)	In-Label	Prot-Inst	Out-Label(s)	Interface/Tunnel-ID	
192.168.48.99 Node	Orig/Transit 10.0.5.185	200099	ISIS-0	200099	toSim185	* * *
3ffe::a14:194 Node	Orig/Transit fe80::62c2:ffff:fe00:0	200294	ISIS-0	200294	toSim194	
3ffe::a14:199 Node	Terminating	200299	ISIS-0			
192.168.48.219 Node	Orig/Transit 10.202.5.194	200319	ISIS-0	200319	toSim194	
192.168.48.221 Node	Orig/Transit 10.0.5.185	200321	ISIS-0	200321	toSim185	
192.168.48.223 Node	Orig/Transit 10.0.5.185	200323	ISIS-0	200323	toSim185	
192.168.48.224 Node	Orig/Transit 10.0.5.185	200324	ISIS-0	200324	toSim185	
192.168.48.226 Node	Orig/Transit 10.0.1.2	200326	ISIS-0	100326	toSim226PCEServer	
192.168.48.184 Node	Orig/Transit 10.0.5.185	200384	ISIS-0	200384	toSim185	
192.168.48.185 Node	Orig/Transit 10.0.5.185	200385	ISIS-0	200385	toSim185	
192.168.48.190 Node	Orig/Transit 10.0.5.185	200390	ISIS-0	200390	toSim185	
192.168.48.194 Node	Orig/Transit 10.202.5.194	200394	ISIS-0	200394	toSim194	
192.168.48.198 Node	Orig/Transit 10.0.9.198	200398	ISIS-0	100398	toSim198	
192.168.48.199 Node	Terminating	200399	ISIS-0			
10.0.9.198 Adjacency	Transit 10.0.9.198	262122	ISIS-0	3	toSim198	
10.202.1.219 Adjacency	Transit 10.202.1.219	262124	ISIS-0	3	toSim219	
10.0.5.185 Adjacency	Transit 10.0.5.185	262133	ISIS-0	3	toSim185	
fe80::62c2:ffff:fe00:0 Adjacency	Transit fe80::62c2:ffff:fe00:0	262134	ISIS-0	3	toSim194	
10.0.1.2 Adjacency	Transit 10.0.1.2	262137	ISIS-0	3	toSim226PCEServer	
10.0.13.184 Adjacency	Transit 10.0.13.184	262138	ISIS-0	3	toSim184	
10.0.2.2 Adjacency	Transit	262139	ISIS-0			

```

10.202.5.194      10.0.2.2      3      toSim226PCEServer202
Adjacency      Transit      262141      ISIS-0
10.202.5.194      3      toSim194
-----
No. of Entries: 22
-----

```

A policy is configured to add the list of prefixes to which the ingress LER Phoenix 199 must auto-create SR-TE LSPs.

```

*A:Phoenix 199>config>router>policy-options# info
-----
    prefix-list "sr-te-level2"
        prefix 192.168.48.198/32 exact
        prefix 192.168.48.221/32 exact
        prefix 192.168.48.223/32 exact
    exit
    policy-statement "sr-te-auto-lsp"
        entry 10
            from
                prefix-list "sr-te-level2"
            exit
            action accept
            exit
        exit
        default-action drop
        exit
    exit
-----

```

An LSP template of type **mesh-p2p-srte** is configured, which uses a path with a loose hop corresponding to anycast SID prefix of the transit routers. The LSP template is then bound to the policy containing the prefix list. See lines marked with an asterisk (*).

```

*A:Phoenix 199>config>router>mpls# info
-----
    cspf-on-loose-hop
    interface "system"
        no shutdown
    exit
    interface "toESS195"
        no shutdown
    exit
    interface "toSim184"
        no shutdown
    exit
    interface "toSim185"
        admin-group "bottom"
        srlg-group "bottom"
        no shutdown
    exit
    interface "toSim194"
        admin-group "bottom"
        srlg-group "bottom"
        no shutdown
    exit
    interface "toSim198"
        no shutdown
    exit
    interface "toSim219"
        no shutdown

```

```

exit
path "loose-anycast-sid"
  hop 1 192.168.48.99 loose
  no shutdown
exit
lsp-template "sr-te-level2-mesh" mesh-p2p-srte
  default-path "loose-anycast-sid"
  max-sr-labels 8 additional-frr-labels 2
  pce-report enable
  no shutdown
exit
auto-lsp lsp-template "sr-te-level2-mesh" policy "sr-te-auto-lsp"
no shutdown
-----

```

One SR-TE LSP is automatically created to each destination matching the prefix in the policy as soon as the router with the router ID matching the address of the prefix appears in the TE database.

The following shows the three SR-TE auto-LSPs created. See lines marked with an asterisk (*).

```

*A:Phoenix 199# show router mpls sr-te-lsp
=====
MPLS SR-TE LSPs (Originating)
=====

```

LSP Name	To	Tun Id	Protect Path	Adm	Opr	
Phoenix-SL-1	192.168.48.223	1	N/A	Up	Up	
Phoenix-SL-2-Profile	192.168.48.223	2	N/A	Up	Up	
Phoenix-SL-3-Profile	192.168.48.223	3	N/A	Up	Up	
Phoenix-SL-4-Profile	192.168.48.223	4	N/A	Up	Up	
Phoenix-SL-1-Profile	192.168.48.223	5	N/A	Up	Up	
Phoenix-SL-2	192.168.48.223	6	N/A	Up	Up	
Phoenix-SL-3	192.168.48.223	7	N/A	Up	Up	
Phoenix-SL-4	192.168.48.223	8	N/A	Up	Up	
sr-te-level2-mesh-192.168.48.198-716803	192.168.48.198	61442	N/A	Up	Up	*
sr-te-level2-mesh-192.168.48.221-716804	192.168.48.221	61443	N/A	Up	Up	*
sr-te-level2-mesh-192.168.48.223-716805	192.168.48.223	61444	N/A	Up	Up	*

```

-----
LSPs : 17
=====

```

The auto-generated name uses the syntax convention *TemplateName-DestIpv4Address-TunnelId*, as described in [Automatic creation of an SR-TE mesh LSP](#). The tunnel ID used in the name is the TTM tunnel ID, not the MPLS LSP tunnel ID. See lines marked with an asterisk (*).

```

*A:Phoenix 199# show router mpls sr-te-lsp "sr-te-level2-mesh-192.168.48.223-716805" detail
=====
MPLS SR-TE LSPs (Originating) (Detail)
=====

```

```

Type : Originating
-----

```

LSP Name	: sr-te-level2-mesh-192.168.48.223-716805		
LSP Type	: MeshP2PSrTe	LSP Tunnel ID	: 61444 *
LSP Index	: 126979	TTM Tunnel Id	: 716805 *
From	: 192.168.48.199	To	: 192.168.48.223
Adm State	: Up	Oper State	: Up

```

LSP Up Time      : 0d 00:02:12      LSP Down Time    : 0d 00:00:00
Transitions      : 3                Path Changes     : 3
Retry Limit      : 0                Retry Timer       : 30 sec
CSPF             : Enabled
Metric           : N/A              Use TE metric     : Disabled
Include Grps     :                  Exclude Grps      :
None
VprnAutoBind     : Enabled
IGP Shortcut     : Enabled          BGP Shortcut      : Enabled
IGP LFA          : Disabled         IGP Rel Metric    : Disabled
BGPTransTun     : Enabled
Oper Metric      : 16777215
PCE Report       : Enabled
PCE Compute      : Disabled        PCE Control       : Disabled
Max SR Labels    : 8                Additional FRR Labels: 2
Path Profile     :
None
Primary(a)       : loose-anycast-sid Up Time           : 0d 00:02:12
Bandwidth        : 0 Mbps
=====

```

The automatically created SR-TE auto-LSPs are also added into the tunnel table to be used by services and shortcut applications. See lines marked with an asterisk (*).

```

*A:Phoenix 199# show router tunnel-table
=====
IPv4 Tunnel Table (Router: Base)
=====
Destination      Owner      Encap TunnelId Pref  Nexthop      Metric
-----
10.0.5.185/32    isis (0)  MPLS  524370      11    10.0.5.185    0
10.0.9.198/32    isis (0)  MPLS  524368      11    10.0.9.198    0
10.0.13.184/32   isis (0)  MPLS  524340      11    10.0.13.184   0
10.202.1.219/32  isis (0)  MPLS  524333      11    10.202.1.219  0
10.202.5.194/32  isis (0)  MPLS  524355      11    10.202.5.194  0
10.0.1.2/32      isis (0)  MPLS  524364      11    11.0.1.2      0
10.0.2.2/32      isis (0)  MPLS  524363      11    11.0.2.2      0
192.168.48.99/32 isis (0)  MPLS  524294      11    10.0.5.185    10
192.168.48.184/32 ldp       MPLS  65605       9     10.0.5.185    20
192.168.48.184/32 isis (0)  MPLS  524341      11    10.0.5.185    20
192.168.48.185/32 ldp       MPLS  65602       9     10.0.5.185    10
192.168.48.185/32 isis (0)  MPLS  524371      11    10.0.5.185    10
192.168.48.190/32 ldp       MPLS  65606       9     10.0.5.185    40
192.168.48.190/32 isis (0)  MPLS  524362      11    10.0.5.185    40
192.168.48.194/32 ldp       MPLS  65577       9     10.202.5.194  10
192.168.48.194/32 isis (0)  MPLS  524331      11    10.202.5.194  10
192.168.48.198/32 sr-te     MPLS  716803      8     192.168.48.99 16777215    *
192.168.48.198/32 ldp       MPLS  65601       9     10.0.9.198    10
192.168.48.198/32 isis (0)  MPLS  524369      11    10.0.9.198    10
192.168.48.219/32 ldp       MPLS  65579       9     10.202.5.194  20
192.168.48.219/32 isis (0)  MPLS  524334      11    10.202.5.194  20
192.168.48.221/32 sr-te     MPLS  716804      8     192.168.48.99 16777215    *
192.168.48.221/32 ldp       MPLS  65607       9     10.0.5.185    30
192.168.48.221/32 isis (0)  MPLS  524358      11    10.0.5.185    30
192.168.48.223/32 sr-te     MPLS  655362      8     10.0.13.184   200
192.168.48.223/32 sr-te     MPLS  655363      8     10.0.13.184   200
192.168.48.223/32 sr-te     MPLS  655364      8     10.0.5.185    40
192.168.48.223/32 sr-te     MPLS  655365      8     10.0.13.184   120
192.168.48.223/32 sr-te     MPLS  655366      8     10.0.5.185    120
192.168.48.223/32 sr-te     MPLS  655367      8     10.0.13.184   120
192.168.48.223/32 sr-te     MPLS  655368      8     10.0.13.184   200
192.168.48.223/32 sr-te     MPLS  655369      8     10.0.5.185    40
192.168.48.223/32 sr-te     MPLS  716805      8     192.168.48.99 16777215    *

```

```

192.168.48.223/32 ldp      MPLS  65603    9      10.0.5.185    20
192.168.48.223/32 isis (0) MPLS  524306   11     10.0.5.185    20
192.168.48.224/32 ldp      MPLS  65604    9      10.0.5.185    30
192.168.48.224/32 isis (0) MPLS  524361   11     10.0.5.185    30
192.168.48.226/32 isis (0) MPLS  524365   11     11.0.1.2      65534

```

```

-----
Flags: B = BGP backup route available
      E = inactive best-external BGP route
=====

```

The details of the path of one of the SR-TE auto-LSPs now show the ERO transiting through the anycast SID of router Houston 185. See lines marked with an asterisk (*).

```

*A:Phoenix 199# show router mpls sr-te-lsp "sr-te-level2-mesh-192.168.48.223-
716805" path detail
=====
MPLS SR-TE LSP sr-te-level2-mesh-192.168.48.223-716805 Path (Detail)
=====
Legend :
  S      - Strict                      L      - Loose
  A-SID  - Adjacency SID              N-SID  - Node SID
  +      - Inherited
=====
SR-TE LSP sr-te-level2-mesh-192.168.48.223-716805 Path loose-anycast-sid
-----
LSP Name          : sr-te-level2-mesh-192.168.48.223-716805
Path LSP ID       : 20480
From              : 192.168.48.199      To              : 192.168.48.223
Admin State       : Up                  Oper State       : Up
Path Name         : loose-anycast-sid   Path Type        : Primary
Path Admin        : Up                  Path Oper        : Up
Path Up Time      : 0d 02:30:28          Path Down Time   : 0d 00:00:00
Retry Limit       : 0                    Retry Timer       : 30 sec
Retry Attempt     : 1                    Next Retry In    : 0 sec
CSPF              : Enabled              Oper CSPF        : Enabled
Bandwidth         : No Reservation        Oper Bandwidth    : 0 Mbps
Hop Limit         : 255                   Oper HopLimit     : 255
Setup Priority     : 7                     Oper Setup Priority : 7
Hold Priority      : 0                     Oper Hold Priority : 0
Inter-area        : N/A
PCE Updt ID       : 0                    PCE Updt State   : None
PCE Upd Fail Code : noError
PCE Report        : Enabled              Oper PCE Report   : Disabled
PCE Control       : Disabled             Oper PCE Control   : Disabled
PCE Compute       : Disabled
Include Groups    :                      Oper Include Groups :
None                                                       None
Exclude Groups   :                      Oper Exclude Groups :
None                                                       None
IGP/TE Metric     : 16777215             Oper Metric       : 16777215
Oper MTU          : 1492                  Path Trans        : 1
Failure Code      : noError
Failure Node      : n/a
Explicit Hops     :
  192.168.48.99(L)
Actual Hops       :
  192.168.48.99 (192.168.48.185) (N-SID)      Record Label     : 200099   *
-> 192.168.48.223 (192.168.48.223) (N-SID)    Record Label     : 200323   *
=====

```

2.2.20 EL on SR-TE LSPs

The router supports the MPLS entropy label on SR-TE LSPs as described in RFC 6790. LSR nodes in a network can load balance labeled packets more granularly than by hashing on the standard label stack. See the *7705 SAR Gen 2 MPLS Guide* for more information.

To allow the EL in the label stack for packets on a SR-TE LSP, the head end router must be able to determine the following:

- Is the far end of the LSP Entropy Label Capability (ELC)?
- Should two additional LSEs (Entropy Label and Entropy Label Indicator - EL/ELI) be added to the SR-TE LSP? This check is required to prevent cases where the additional LSEs may cause the maximum SID depth to be exceeded.

The EL can be inserted on packets from a service that is configured to use EL if both of these criteria are satisfied.

Announcing the ELC is supported by the OSPF or IS-IS routing protocol. However, processing the ELC signaling is not supported for OSPF or IS-IS segment-routed tunnels. That is, the router does not take into account the entropy label capability received in advertisements from nodes in the IS-IS or OSPF domain when determining if the far end of an SR-TE LSP is capable of receiving and processing packets containing the entropy label. ELC is therefore determined by configuring the **override-tunnel-elc** command, either under the IS-IS or OSPF IGP configuration, or under the SR-TE LSP itself, as described in the following paragraphs. In addition, use the following command to instruct MPLS that EL/ELI can be inserted on SR-TE LSPs.

```
configure router mpls entropy-label sr-te
```

This command applies to all SR-TE LSPs originating on the router. Note that this global configuration can be overridden on an LSP by LSP basis using the following command.

```
configure router mpls lsp entropy-label
```

If the path of the SR-TE LSP is computed by the local CSPF or IP-to-label translation, the head-end router assumes ELC if either the following commands is configured.

```
configure router isis entropy-label override-tunnel-elc
configure router ospf entropy-label override-tunnel-elc
```

However, this case requires that the far-end node SID of the LSP is advertised within the same domain as the head end. This allows the head-end router to know the association between the far-end IP address and the SID of the node for which to insert the EL.

When some types of SR-TE LSP paths are specified as a list of SID labels, the head-end LER cannot derive the ELC of the SR-TE LSP from the IGP. It therefore needs to be explicitly configured for each LSP. This applies to the following cases:

- **for SR-TE LSPs where the primary or secondary path hops consist of static SID labels**

The SID labels are configured under the following context.

```
configure router mpls path hop
```

In this case, use the following command to configure the ELC.

```
configure router mpls lsp override-tunnel-elc
```

2.3 Segment routing policies

The concept of an SR policy is described in RFC 9256. An SR policy specifies a source-routed path from a head-end router to a network endpoint, and the traffic flows that are steered to that source-routed path. An SR policy intended for use by a specific head-end router can be statically configured on that router or advertised to it in the form of a BGP route.

The following terms describe the structure of an SR policy and the relationship between one policy and another.

- **SR policy**

This policy is identified by the tuple of <headend, color, endpoint>. Each SR policy is associated with a set of one or more candidate paths, one of which is selected to implement the SR policy and is installed in the data plane. Certain properties of the SR policy come from the currently selected path, such as binding SID, segment lists, and so on.

- **endpoint**

This is the far-end router that is the destination of the source-routed path. The endpoint may be null (all-zero IP address) if no specific far-end router is targeted by the policy.

- **color**

This property of an SR policy determines the sets of traffic flows that are steered by the policy.

- **path**

This is a set of one or more segment lists that are explicitly or statically configured or dynamically signaled. If a path becomes active, traffic matching the SR policy is load-balanced across the segment lists of the path in an equal, unequal, or weighted distribution. Each path is associated with:

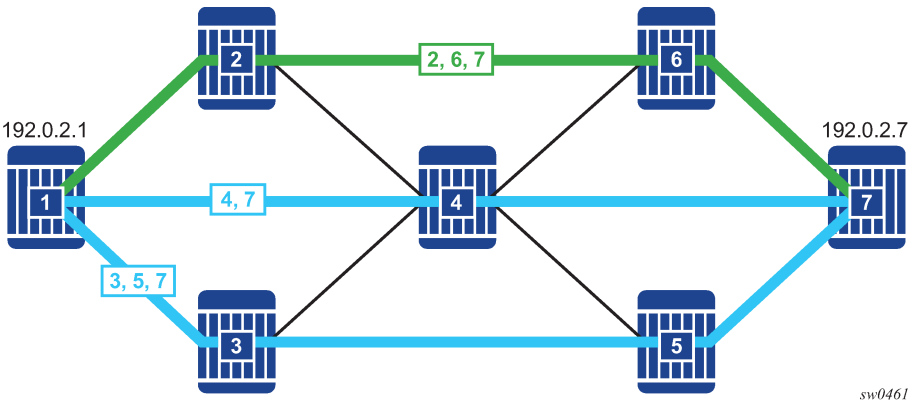
- a protocol origin (BGP or static)
- a preference value
- a binding SID value
- a validation state (valid or invalid)

- **BSID**

The binding SID (BSID) value opaquely represents an SR policy (or more specifically, its selected path) to upstream routers. BSIDs provide isolation or decoupling between different source-routed domains and improve overall network scalability. Usually, all candidate paths of an SR policy are assigned the same BSID.

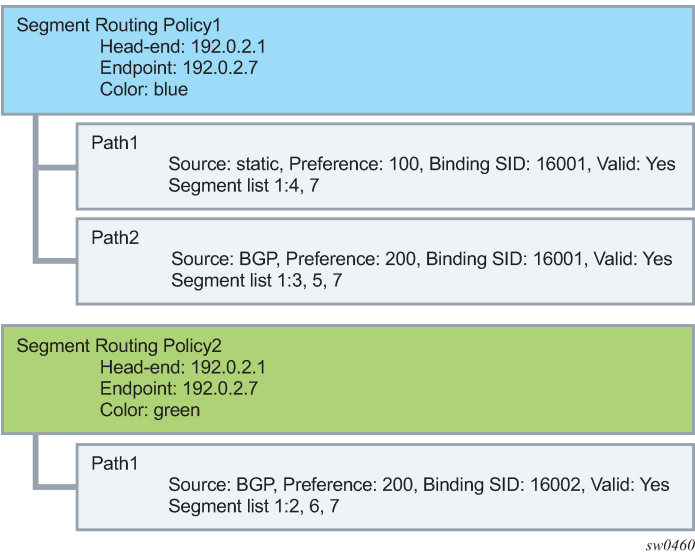
These concepts are illustrated in the following example. Suppose there is a network of seven nodes, as shown in the following figure, and there are two classes of traffic (blue and green) to be transported between node 1 and node 7. There is an SR policy for the blue traffic between node 1 and node 7 and another SR policy for the green traffic between these two nodes.

Figure 25: Network example with two SR policies



The two SR policies that are involved in this example and the associated relationships are depicted in [Figure 26: Relationship between SR policies and paths](#).

Figure 26: Relationship between SR policies and paths



2.3.1 Statically-configured segment routing policies

An SR policy is statically configured on the router using one of the supported management interfaces. In the Nokia data model, static policies are configured under **config>router>segment-routing>sr-policies**.

There are two types of static policies: local and non-local. A static policy is local when its **head-end** parameter is configured with the value *local*. This means that the policy is intended for use by the router where the static policy is configured. Local static policies are imported into the local segment routing database for further processing. If the local segment routing database chooses a local static policy as the best path for a specific (color, endpoint) combination, then the associated path and its segment lists are installed into the tunnel table (for next-hop resolution) and as a BSID-indexed MPLS label entry.

A static policy is non-local when its **head-end** parameter is set to any IPv4 address (even an IPv4 address that is associated with the local router, which is a configuration that should generally be avoided). A non-local policy is intended for use by a different router than the one where the policy is configured. Non-local policies are not installed in the local segment routing database and do not affect the forwarding state of the router where they are configured. To advertise non-local policies to the target router, either directly (over a single BGP session) or indirectly (using other intermediate routers, such as BGP route reflectors), the static non-local policies must be imported into the BGP RIB and then re-advertised as BGP routes. To import static non-local policies into BGP, the user must configure the **sr-policy-import** command under **config>router>bgp**. To advertise BGP routes containing SR policies, the user must add the **sr-policy-ipv4** or the **sr-policy-ipv6** family to the configuration of a BGP neighbor or group (or the entire Base router BGP instance) so that the capability is negotiated with other routers.

Local and non-local static policies have the same configurable attributes. The function and rules associated with each attribute are:

- **shutdown**

This command is used to administratively enable or disable the static policy.

- **binding-sid**

This command is used to associate a binding SID with the static policy in the form of an MPLS label in the range of 32 to 1048575. This is a mandatory parameter. The binding SID must be an available label in the reserved label block associated with SR policies, otherwise the policy cannot be activated.

- **color**

This command is used to associate a color with the static policy. This is a mandatory parameter.

- **distinguisher**

This command is used to uniquely identify a non-local static policy when it is a re-advertised as a BGP route. The value is copied into the BGP NLRI field. A unique distinguisher ensures that BGP does not suppress BGP routes for the same (color, endpoint), but is targeted to different head-end routers. This is mandatory for non-local policies but optional in local policies.

- **endpoint**

This command is used to identify the endpoint IPv4 or IPv6 address associated with the static policy. A value of 0.0.0.0 or ::0 is permitted and is interpreted as a null endpoint. This is a mandatory parameter.



Note: When a non-local SR policy with either an IPv4 or IPv6 endpoint is selected for advertisement, the **head-end** command supports an IPv4 address only. This is converted into an IPv4-address-specific RT extended community (0x4102) in the advertised route in the BGP Update message.

- **head-end**

This command is used to identify the router that is the targeted node for installing the policy. This is a mandatory parameter. The **local** parameter must be used when the target is the local router itself. Otherwise, any valid IPv4 address is allowed, and the policy is considered non-local. When a non-local static policy is re-advertised as a BGP route, the configured head-end address is embedded in an IPv4-address-specific route-target extended community that is automatically added to the BGP route.

- **preference**

This command is used to indicate the degree of preference of the policy if the local segment routing database has other policies (static or BGP) for the same (color, endpoint) combination. For a path to

be selected as the active path for a (color, endpoint) combination, it must have the highest preference value amongst all the candidate paths.

The following are configuration rules related to the previously described attributes.

- Every static local policy must have a unique combination of color, endpoint, and preference.
- Every static non-local policy must have a unique distinguisher.

Each static policy (local and non-local) must include at least one segment list containing at least one segment in its configuration. Each static-policy can have up to 32 segment lists, each containing up to 11 segments. Each segment list can be assigned a weight to influence the share of traffic that it carries compared to other segment lists of the same policy. The default weight is 1.

The segment routing policy draft standard allows a segment list to be configured (and signaled) with a mix of different segment types. When the head-end router attempts to install such a segment routing policy, it must resolve all of the segments into a stack of MPLS labels. In SR OS, this complexity is avoided by requiring that all configured and signaled segments must already be provided in the form of MPLS label values. As described in the draft, this means that only type-1 segments are supported.

2.3.2 BGP-signaled SR policies

The Base router BGP instance is configured to send and receive BGP routes containing SR policies. To exchange routes belonging to the (AFI=1, SAFI=73) or (AFI=2, SAFI=73) address family with a specific Base router BGP neighbor, the family configuration that applies to that neighbor must include the **sr-policy-ipv4** or the **sr-policy-ipv6** keyword respectively.

When BGP receives an **sr-policy-ipv4** route (AFI=1, SAFI=73) or a **sr-policy-ipv6 route** (AFI=2, SAFI=73) from a peer, it runs its standard BGP best path selection algorithm to choose the best path for each NLRI combination of distinguisher, endpoint, and color. If the best path is targeted to this router as the head end, BGP extracts the SR policy details into the local SR database. A BGP SR policy route is deemed to be targeted to this router as the head end if either:

- it has no route-target extended community and a NO-ADVERTISE standard community
- it has an IPv4 address-specific route-target extended community with an IPv4 address matching the system IPv4 address of this router

An **sr-policy-ipv4** or a **sr-policy-ipv6** route can be received from either an IBGP or EBGP peer but it is never propagated to an EBGP peer. An **sr-policy-ipv4** or a **sr-policy-ipv6** route can be reflected to route reflector clients if this is allowed (a NO_ADVERTISE community is not attached) and the router does not consider itself the head end of the policy.



Note: A BGP SR policy route is considered malformed if it does not have at least one segment list TLV with at least one segment TLV, which triggers error-handling procedures such as session reset or treat-as withdraw.

2.3.3 Segment routing policy path selection and tie-breaking

Segment Routing policies (static and BGP) for which the local router is the head end are processed by the local segment routing database. For each (color, endpoint) combination, the database validates each candidate path and chooses one to be the active path. The steps of this process are described in the Segment Routing policy validation and selection process.

1. Is the path missing a binding SID in the form of an MPLS label?

- Yes: this path is invalid and cannot be used.
 - No: go to the next step.
2. Does the path have any segment list that contains a segment type not equal to 1 (that is, an MPLS label)?
- Yes: this path is invalid and cannot be used.
 - No: go to the next step.

node-SID

3. Are all segment lists of the path invalid?
A segment list is invalid if it is empty, if the first SID cannot be resolved to a set of one or more next-hops, or if the weight is 0.
- Yes: this path is invalid and cannot be used.
 - No: go to the next step.

At this step, the router attempts to resolve the first segment of each segment list to a set of one or more next-hops and outgoing labels. It does so by looking for a matching SID in the segment routing module, which must correspond to one of the following:

- SR IS-IS or SR-OSPF node SID
- SR IS-IS or SR-OSPF adjacency SID
- SR IS-IS or SR-OSPF adjacency set SID (parallel or non-parallel set)



Note: The label value in the first segment of the segment list is matched against ILM label values that the local router has assigned to the node SIDs, adjacency SIDs, and adjacency set SIDs. The matched ILM entry may not program a swap to the same label value encoded in the segment routing policy; for example, in the case of an adjacency SID or of a node SID reachable through a next hop using a different SRGB base.

4. Is the binding SID an available label in the reserved-label-block range?
- Yes: go to the next step.
 - No: this path is invalid and cannot be used.
5. Is there another path that has reached this step that has a higher preference value?
- Yes: this path loses the tie-break and cannot be used.
 - No: go to the next step.
6. Is there a static path?
- Yes: select the static path as the active path because the protocol-origin value associated with static paths (30) is higher than the protocol-origin value associated with BGP learned paths (20).
 - No: go to the next step.
7. Is there a BGP path with a lower originator value?
The originator is a 160-bit numerical value formed by the concatenation of a 32-bit ASN and a 128-bit peer address (with IPv4 addresses encoded in the lowest 32 bits).
- Yes: this path loses the tie-break and cannot be used.
8. Is there another BGP path with a higher distinguisher value?
- Yes: select the BGP path with the highest distinguisher value.

2.3.4 Resolving BGP routes to segment routing policy tunnels

When a statically configured or BGP-signaled segment routing policy is selected to be the active path for a (color, endpoint) combination, the corresponding path and its segment lists are programmed into the tunnel table of the router. An IPv4 tunnel of type **sr-policy** (where the **endpoint** parameter is an IPv4 address) is programmed into the IPv4 tunnel table (TTMv4). Similarly, an IPv6 tunnel of type **sr-policy** (where the **endpoint** parameter is an IPv6 address) is programmed into the IPv6 tunnel table (TTMv6). The resulting tunnel entries can be used to resolve the following types of BGP routes:

- Unlabeled IPv4 routes
- Unlabeled IPv6 routes
- Label-unicast IPv4 routes
- Label-unicast IPv6 (6PE) routes
- VPN IPv4 and IPv6 routes
- EVPN routes

Specifically, an IPv4 tunnel of type **sr-policy** can be used to resolve:

- an IPv4 or the IPv4-mapped IPv6 next hop of the following route families:
ipv4, ipv6, vpn-ipv4, vpn-ipv6, label-ipv4, label-ipv6, evpn
- the IPv6 next hop of the following route families:
ipv6, label-ipv4, and label-ipv6 (SR policy with **endpoint** = 0.0.0.0 only).

An IPv6 tunnel of type **sr-policy** can be used to resolve:

- the IPv6 next hop of the following route families:
ipv4, ipv6, vpn-ipv4, vpn-ipv6, label-ipv4, label-ipv6, evpn
- the IPv4 next hop of the following route families:
ipv4 and label-ipv4 (SR policy with **endpoint** = 0::0 only).
- the IPv4-mapped IPv6 next hop of the following route families:
label-ipv6 (SR policy with **endpoint** = 0::0 only).

2.3.4.1 Resolving unlabeled IPv4 BGP routes to segment routing policy tunnels

For an unlabeled IPv4 BGP route to be resolved by an SR policy:

- A color-extended community must be attached to the IPv4 route.
- The base instance BGP next-hop-resolution configuration of **shortcut-tunnel>family ipv4** must allow SR policy tunnels.



Note: Contrary to *draft-filsfils-segment-routing-05*, BGP only resolves a route with multiple color-extended communities to an SR policy using the color-extended community with the highest value.

For example, to resolve an IPv4 route with a color-extended community (value C) and BGP next-hop address N under these conditions, the router performs the following steps:

1. If there is an SR policy in the TTMv4 for the following are true, then use this tunnel to resolve the BGP next hop:
 - end-point = BGP next-hop address
 - color = *Cn*
2. If no SR policy is found in the previous step and the *Cn* color-extended community has its color-only (CO) bits set to 01 or 10, then search for an SR policy in the TTMv4 for which the following are true:
 - endpoint = null (0.0.0.0)
 - color = *Cn*
 If there is such a policy, use it to resolve the BGP next hop.
3. If no SR policy is found in the previous steps and the *Cn* color-extended community has its CO bits set to 01 or 10, then search for an SR policy in the TTMv6 for which the following are true.
 - endpoint = null (0::0)
 - color = *Cn*
 If there is such a policy, use it to resolve the BGP next hop.
4. If no SR policy is found in the previous steps but there is a non-SR policy tunnel in the TTMv4 that is allowed by the resolution options and for which endpoint = BGP next-hop address (and for which the admin tag meets the admin-tag-policy requirements applied to the BGP route, if applicable), then use this tunnel to resolve the BGP next hop if it has the highest TTM preference.
5. Otherwise, fall back to IGP, unless the **disallow-igp** option is configured.

2.3.4.2 Resolving unlabeled IPv6 BGP routes to segment routing policy tunnels

For an unlabeled IPv6 BGP route to be resolved by an SR policy:

- A color-extended community must be attached to the IPv6 route.
- The base instance BGP next-hop-resolution configuration of **shortcut-tunnel>family ipv6** must allow SR policy tunnels.



Note:

- Contrary to *draft-filsfils-segment-routing-05*, BGP only resolves a route with multiple color-extended communities to an SR policy using the color-extended community with the highest value.
- For AFI2/SAFI1 routes, an IPv6 explicit null label is pushed at the bottom of the stack if the policy endpoint is IPv4.

For example, to resolve an IPv6 route with a color-extended community (value *C*) and BGP next-hop address *N* under these conditions, the router performs the following steps:

1. If there is an SR policy in the TTMv6 for which the following are true, then use this tunnel to resolve the BGP next hop.
 - endpoint = the BGP next-hop address
 - color = *Cn*
2. If no SR policy is found in the previous step and the *Cn* color-extended community has its CO bits set to 01 or 10, then search for an SR policy in the TTMv6 for which the following are true:

- endpoint = null (0::0)
- color = *Cn*

If there is such a policy, use it to resolve the BGP next hop.

3. If no SR policy is found in the previous steps, and the *Cn* color-extended community has its CO bits set to 01 or 10 and there is an SR policy in the TTMv4 for which the following are true, then use this tunnel to resolve the BGP next hop.
 - endpoint = null (0.0.0.0)
 - color = *Cn*
4. If no SR policy is found in the previous steps but there is a non-SR policy tunnel in the TTMv6 that is allowed by the resolution options and for which endpoint = BGP next-hop address (and for which the admin-tag meets the admin-tag-policy requirements applied to the BGP route, if applicable), then use this tunnel to resolve the BGP next hop if it has the highest TTM preference.
5. Otherwise, fall back to IGP, unless the **disallow-igp** option is configured.

2.3.4.3 Resolving label-IPv4 BGP routes to segment routing policy tunnels

For a label-unicast IPv4 BGP route to be resolved by an SR policy:

- A color-extended community must be attached to the label-IPv4 route.
- The base instance BGP next-hop-resolution configuration of **labeled-routes>transport-tunnel>family label-ipv4** must allow SR policy tunnels.



Note: Contrary to *draft-filsfils-segment-routing-05*, BGP only resolves a route with multiple color-extended communities to an SR policy using the color-extended community with the highest value.

For example, to resolve a label-IPv4 route with a color-extended community (value *C*) and BGP next-hop address *N*, the router performs the following steps:

1. If there is an interface route that can resolve the BGP next hop, then use the direct route.
2. If the **allow-static** command is configured and there is a static route that can resolve the BGP next hop, then use the static route.
3. If there is no interface route or static route available or allowed to resolve the BGP next hop and the next hop uses IPv4, then:
 - a. Look for an SR policy in the TTMv4 for which the following are true:
 - end-point = BGP next-hop address
 - color = *Cn*
 If there is such an SR policy, use it to resolve the BGP next hop. If the selected SR policy has any segment list with more than {11 - **max-sr-frr-labels** under the IGPs} labels or segments, then the label-IPv4 route is unresolved.
 - b. If no SR policy is found in the previous steps and the *Cn* color-extended community has its CO bits set to 01 or 10, then search for an SR policy in the TTMv4 for which the following are true:
 - endpoint = null (0.0.0.0)
 - color = *Cn*

If there is such a policy, use it to resolve the BGP next hop. If the selected SR policy has any segment list with more than {11- **max-sr-frr-labels** under the IGPs} labels or segments, then the label-IPv4 route is unresolved.

- c. If no SR policy is found in the previous steps and the *Cn* color-extended community has its CO bits set to 01 or 10, then search for an SR policy in the TTMv6 for the following are true:

- endpoint = null (0::0)
- color = *Cn*

If there is such a policy, use it to resolve the BGP next hop. If the selected SR policy has any segment list with more than {11- **max-sr-frr-labels** under the IGPs} labels or segments, then the label-IPv4 route is unresolved.

4. If there is no interface route or static route that is available or allowed to resolve the BGP next hop and the next hop uses IPv6, then:

- a. Search for an SR policy in the TTMv6 for which the following are true:

- end-point = BGP next-hop address
- color = *Cn*

If there is such an SR policy, use it to resolve the BGP next hop. If the selected SR policy has any segment list with more than {11- **max-sr-frr-labels** under the IGPs} labels or segments, then the label-IPv4 route is unresolved.

- b. If no SR policy is found in the previous steps and the *Cn* color-extended community has its CO bits set to 01 or 10, then search for an SR policy in the TTMv6 for which the following are true:

- endpoint = null (0::0)
- color = *Cn*

If there is such a policy, use it to resolve the BGP next hop. If the selected SR policy has any segment list with more than {11- **max-sr-frr-labels** under the IGPs} labels or segments, then the label-IPv4 route is unresolved.

- c. If no SR policy is found in the previous steps and the *Cn* color-extended community has its CO bits set to 01 or 10, then search for an SR policy in the TTMv4 for which the following are true:

- endpoint = null (0.0.0.0)
- color = *Cn*

If there is such a policy, use it to resolve the BGP next hop. If the selected SR policy has any segment list with more than {11- **max-sr-frr-labels** under the IGPs} labels or segments, then the label-IPv4 route is unresolved.

5. If no SR policy is found in the previous steps but there is a non-SR policy tunnel in the TTMv4 (the next hop uses IPv4) or TTMv6 (the next hop uses IPv6) that is allowed by the resolution options and for which endpoint = BGP next-hop address (and for which the admin tag meets the admin-tag-policy requirements applied to the BGP route, if applicable), then use this tunnel to resolve the BGP next hop if it has the highest TTM preference.

2.3.4.4 Resolving label-IPv6 BGP routes to segment routing policy tunnels

For a label-unicast IPv6 BGP route to be resolved by an SR policy:

- A color-extended community must be attached to the label-IPv6 route.

- The base instance BGP next-hop-resolution configuration of the **labeled-routes>transport-tunnel>family label-ipv6** command must allow SR policy tunnels.



Note: Contrary to *draft-filsfils-segment-routing-05*, BGP only resolves a route with multiple color-extended communities to an SR policy using the color-extended community with the highest value.

For example, to resolve a label-IPv6 route with a color-extended community (value *C*) and BGP next-hop address *N*, the router performs the following steps:

1. If there is an interface route that can resolve the BGP next hop, then use the direct route.
2. If the **allow-static** command is configured and there is a static route that can resolve the BGP next hop, then use the static route.
3. If there is no interface route or static route available or allowed to resolve the BGP next hop and the next hop uses IPv6 then:
 - a. Look for an SR policy in the TTMv6 for which the following are true:
 - end-point = BGP next-hop address
 - color = *Cn*
 If there is such an SR policy, use it to resolve the BGP next hop.
 - b. If no SR policy is found in the previous steps and the *Cn* color-extended community has its CO bits set to 01 or 10, then search for an SR policy in the TTMv6 for which the following are true:
 - endpoint = null (0::0)
 - color = *Cn*
 If there is such a policy, use it to resolve the BGP next hop.
 - c. If no SR policy is found in the previous steps and the *Cn* color-extended community has its CO bits set to 01 or 10 then search for an SR policy in the TTMv4 for which the following are true:
 - endpoint = null (0.0.0.0)
 - color = *Cn*
 If there is such a policy, use it to resolve the BGP next hop.
4. If there is no interface route or static route that is available or allowed to resolve the BGP next hop and the next hop uses IPv4-mapped-IPv6, then:
 - a. Look for an SR policy in the TTMv4 for which the following are true:
 - end-point = BGP next-hop address
 - color = *Cn*
 If there is such an SR policy then use it to resolve the BGP next hop.
 - b. If no SR policy is found in the previous steps and the *Cn* color-extended community has its CO bits set to 01 or 10, then search for an SR policy in the TTMv4 for which the following are true:
 - endpoint = null (0.0.0.0)
 - color = *Cn*
 If there is such a policy, use it to resolve the BGP next hop.
 - c. If no SR policy is found in the previous steps and the *Cn* color-extended community has its CO bits set to 01 or 10, then search for an SR policy in the TTMv6 for which the following are true:

- endpoint = null (0::0)
- color = C_n

If there is such a policy, use it to resolve the BGP next hop.

5. If no SR policy is found in the previous steps but there is a non-SR-policy tunnel in the TTMv6 (the next hop uses IPv6) or in the TTMv4 (the next hop uses IPv4-mapped IPv6) that is allowed by the resolution options and for which endpoint = BGP next-hop address (and for which the admin-tag meets the admin-tag-policy requirements applied to the BGP route, if applicable), then use this tunnel to resolve the BGP next hop if it has the highest TTM preference.

2.3.4.5 Resolving EVPN-MPLS routes to segment routing policy tunnels

The next-hop resolution for all EVPN-VXLAN routes and for EVPN-MPLS routes without a color-extended community is unchanged by this feature.

When the resolution options associated with the **auto-bind-tunnel** configuration of an EVPN-MPLS service (VPLS, B-VPLS, R-VPLS, or Epipe) allow SR policy tunnels from the TTM, the next-hop resolution of EVPN-MPLS routes (RT-1 per-EVI, RT-2, RT-3 and RT-5) with one or more color-extended communities (C_1 , C_2 , .. C_n , where C_n is the highest value) is based on the following rules.



Note: Contrary to *draft-filsfils-segment-routing-05*, BGP only resolves a route with multiple color-extended communities to an SR policy using the color-extended community with the highest value.

1. If the next hop uses IPv6 and there is an SR policy in the TTMv6 for which the following are true, then use this tunnel to resolve the BGP next hop.
 - end-point = BGP next-hop address
 - color = C_n
2. Otherwise, if the next hop uses IPv4 or IPv4-mapped IPv6 and there is an SR policy in the TTMv4 for which the following are true, then use this tunnel to resolve the BGP next hop.
 - end-point = BGP next-hop address (or the IPv4 address extracted from the IPv4-mapped IPv6 BGP next-hop address)
 - color = C_n
3. If no SR policy is found in the previous steps but there is a non-SR policy tunnel in the TTMv4 (the next hop uses IPv4 or IPv4-mapped IPv6) or TTMv6 (the next hop uses IPv6) that is allowed by the resolution options and for which endpoint = BGP next-hop address, then use this tunnel to resolve the BGP next hop if it has the highest TTM preference.

2.3.4.6 VPRN auto-bind-tunnel using segment routing policy tunnels

When the resolution options associated with the **auto-bind-tunnel** configuration of VPRN service allow SR policy tunnels from the TTM, next-hop resolution of VPN-IPv4 and VPN-IPv6 routes that are imported into the VPRN and have one or more color-extended communities (C_1 , C_2 , .. C_n , where C_n is the highest value) is based on the following rules.



Note: Contrary to *draft-filsfils-segment-routing-05*, BGP only resolves a route with multiple color-extended communities to an SR policy using the color-extended community with the highest value.

1. If the next hop uses IPv6 and there is an SR policy in the TTMv6 for which the following are true, then use this tunnel to resolve the BGP next hop.
 - end-point = BGP next-hop address
 - color = C_n
2. Otherwise, if the next hop uses IPv4 or IPv4-mapped IPv6 and there is an SR policy in the TTMv4 for which the following are true, then use this tunnel to resolve the BGP next hop.
 - , end-point = BGP next-hop address (or the IPv4 address extracted from the IPv4-mapped IPv6 BGP next-hop address in the case of VPN-IPv6 routes)
 - color = C_n
3. If no SR policy is found in the previous steps but there is a non-SR policy tunnel in the TTMv4 (the next hop uses IPv4 or IPv4-mapped IPv6) or TTMv6 (the next hop uses IPv6) that is allowed by the resolution options and for which endpoint = BGP next-hop address, then use this tunnel to resolve the BGP next hop if it has the highest TTM preference.

2.3.5 Traffic statistics for segment routing policies

SR policies provide the ability to collect statistics for ingress and egress traffic. In both cases, traffic statistics are collected without any forwarding class or QoS distinction.

Traffic statistics collection is enabled as follows:

- **configure router segment-routing sr-policies ingress-statistics**

Ingress traffic collection only applies to **binding-sid** SR policies as the statistic index is attached to the ILM entry for that label. The traffic statistics provide traffic for all the instances that share the binding SID. The statistic index is released and statistics are lost when ingress traffic statistics are disabled for that binding SID or when the last instance of a policy using that label is removed from the database.

- **configure router segment-routing sr-policies egress-statistics**

Egress traffic statistics are collected globally for all policies at the same time. Both static and signaled policies are subject to traffic statistics collection. Statistic indexes are allocated per segment list, which allows for a fine grain monitoring of traffic evolution, but are only allocated at the time the segment list is effectively programmed. The system allocates up to 32 statistic indexes across all the instances of a policy.

If an instance of a policy is deprogrammed and a more preferred instance is programmed, the system behaves as follows:

- If the segment list IDs of the preferred instance are different from any of the segment list IDs of any previously programmed instance, the system allocates new statistic indexes. While that condition holds, the statistics associated with a segment list of an instance strictly reflect the traffic that used that segment list in that instance.
- If some of the segment list IDs of the preferred instance are equal to any of the segment list IDs of any previously programmed instance, the system reuses the indexes of the preferred instance and keeps the associated counter value and increment. In this case, the traffic statistics provided per segment list not only reflect the traffic that used that segment list in that instance, but incorporates counter values of at least another segment list in another instance of that policy.

In all cases, the aggregate values provided across all instances truly reflect traffic over the various instances of the policy.

Statistic indexes are not released at deprogramming time. They are, however, released when all the instances of a policy are removed from the database or when the **egress-statistics** command is disabled.

When the **egress-statistics** command is enabled, the user can configure rate computation on egress. The traffic rate is determined by an accounting policy configuration that uses the **combined-sr-policy-egress** command option and then references the accounting policy in the following context.

```
configure router segment-routing sr-policies egress-statistics
```

The minimum collection interval is 5 minutes. Rate statistics are determined per segment list and accessible using the **show snmp** command as well as via YANG or NETCONF.

3 Standards and protocol support

**Note:**

The information provided in this chapter is subject to change without notice and may not apply to all platforms.

Nokia assumes no responsibility for inaccuracies.

3.1 Bidirectional Forwarding Detection (BFD)

RFC 5880, *Bidirectional Forwarding Detection (BFD)*

RFC 5881, *Bidirectional Forwarding Detection (BFD) IPv4 and IPv6 (Single Hop)*

RFC 5882, *Generic Application of Bidirectional Forwarding Detection (BFD)*

3.2 Border Gateway Protocol (BGP)

draft-hares-idr-update-attr-low-bits-fix-01, *Update Attribute Flag Low Bits Clarification*

draft-ietf-idr-add-paths-guidelines-08, *Best Practices for Advertisement of Multiple Paths in IBGP*

draft-ietf-idr-best-external-03, *Advertisement of the best external route in BGP*

draft-ietf-idr-bgp-gr-notification-01, *Notification Message support for BGP Graceful Restart*

draft-ietf-idr-bgp-optimal-route-reflection-10, *BGP Optimal Route Reflection (BGP-ORR)*

draft-ietf-idr-error-handling-03, *Revised Error Handling for BGP UPDATE Messages*

draft-ietf-idr-link-bandwidth-03, *BGP Link Bandwidth Extended Community*

RFC 1772, *Application of the Border Gateway Protocol in the Internet*

RFC 1997, *BGP Communities Attribute*

RFC 2385, *Protection of BGP Sessions via the TCP MD5 Signature Option*

RFC 2439, *BGP Route Flap Damping*

RFC 2545, *Use of BGP-4 Multiprotocol Extensions for IPv6 Inter-Domain Routing*

RFC 2858, *Multiprotocol Extensions for BGP-4*

RFC 2918, *Route Refresh Capability for BGP-4*

RFC 4271, *A Border Gateway Protocol 4 (BGP-4)*

RFC 4360, *BGP Extended Communities Attribute*

RFC 4364, *BGP/MPLS IP Virtual Private Networks (VPNs)*

RFC 4456, *BGP Route Reflection: An Alternative to Full Mesh Internal BGP (IBGP)*

RFC 4486, *Subcodes for BGP Cease Notification Message*

RFC 4659, *BGP/MPLS IP Virtual Private Network (VPN) Extension for IPv6 VPN*

RFC 4684, *Constrained Route Distribution for Border Gateway Protocol/MultiProtocol Label Switching (BGP/MPLS) Internet Protocol (IP) Virtual Private Networks (VPNs)*

RFC 4724, *Graceful Restart Mechanism for BGP – helper mode*

RFC 4760, *Multiprotocol Extensions for BGP-4*

RFC 4798, *Connecting IPv6 Islands over IPv4 MPLS Using IPv6 Provider Edge Routers (6PE)*

RFC 5004, *Avoid BGP Best Path Transitions from One External to Another*

RFC 5065, *Autonomous System Confederations for BGP*

RFC 5291, *Outbound Route Filtering Capability for BGP-4*

RFC 5396, *Textual Representation of Autonomous System (AS) Numbers – asplain*

RFC 5492, *Capabilities Advertisement with BGP-4*

RFC 5668, *4-Octet AS Specific BGP Extended Community*

RFC 6286, *Autonomous-System-Wide Unique BGP Identifier for BGP-4*

RFC 6368, *Internal BGP as the Provider/Customer Edge Protocol for BGP/MPLS IP Virtual Private Networks (VPNs)*

RFC 6793, *BGP Support for Four-Octet Autonomous System (AS) Number Space*

RFC 6810, *The Resource Public Key Infrastructure (RPKI) to Router Protocol*

RFC 6811, *Prefix Origin Validation*

RFC 6996, *Autonomous System (AS) Reservation for Private Use*

RFC 7311, *The Accumulated IGP Metric Attribute for BGP*

RFC 7606, *Revised Error Handling for BGP UPDATE Messages*

RFC 7607, *Codification of AS 0 Processing*

RFC 7752, *North-Bound Distribution of Link-State and Traffic Engineering (TE) Information Using BGP*

RFC 7911, *Advertisement of Multiple Paths in BGP*

RFC 7999, *BLACKHOLE Community*

RFC 8092, *BGP Large Communities Attribute*

RFC 8097, *BGP Prefix Origin Validation State Extended Community*

RFC 8212, *Default External BGP (EBGP) Route Propagation Behavior without Policies*

RFC 8277, *Using BGP to Bind MPLS Labels to Address Prefixes*

RFC 9294, *Application-Specific Link Attributes Advertisement Using the Border Gateway Protocol - Link State (BGP LS)*

RFC 9494, *Long-Lived Graceful Restart for BGP*

3.3 Bridging and management

IEEE 802.1AB, *Station and Media Access Control Connectivity Discovery*

IEEE 802.1ad, *Provider Bridges*

IEEE 802.1AX, *Link Aggregation*

IEEE 802.1D, *MAC Bridges*
IEEE 802.1p, *Traffic Class Expediting*
IEEE 802.1Q, *Virtual LANs*
IEEE 802.1s, *Multiple Spanning Trees*
IEEE 802.1w, *Rapid Reconfiguration of Spanning Tree*

3.4 Certificate management

RFC 4210, *Internet X.509 Public Key Infrastructure Certificate Management Protocol (CMP)*
RFC 4211, *Internet X.509 Public Key Infrastructure Certificate Request Message Format (CRMF)*
RFC 5280, *Internet X.509 Public Key Infrastructure Certificate and Certificate Revocation List (CRL) Profile*
RFC 6712, *Internet X.509 Public Key Infrastructure -- HTTP Transfer for the Certificate Management Protocol (CMP)*
RFC 7030, *Enrollment over Secure Transport*
RFC 7468, *Textual Encodings of PKIX, PKCS, and CMS Structures*

3.5 Ethernet

IEEE 802.3x, *Ethernet Flow Control*

3.6 Ethernet VPN (EVPN)

draft-ietf-bess-evpn-ipvpn-interworking-15, *EVPN Interworking with IPVPN*
draft-ietf-bess-evpn-l3mh-proto-00, *EVPN Multi-Homing support for L3 services*
draft-rbickhart-evpn-ip-mac-proxy-adv-04, *Proxy MAC-IP Advertisement in EVPN*
RFC 7432, *BGP MPLS-Based Ethernet VPN*
RFC 8214, *Virtual Private Wire Service Support in Ethernet VPN*
RFC 8365, *A Network Virtualization Overlay Solution Using Ethernet VPN (EVPN)*
RFC 8560, *Seamless Integration of Ethernet VPN (EVPN) with Virtual Private LAN Service (VPLS) and Their Provider Backbone Bridge (PBB) Equivalents*
RFC 9047, *Propagation of ARP/ND Flags in an Ethernet Virtual Private Network (EVPN)*
RFC 9135, *Integrated Routing and Bridging in Ethernet VPN (EVPN)*
RFC 9136, *IP Prefix Advertisement in Ethernet VPN (EVPN)*
RFC 9161, *Operational Aspects of Proxy ARP/ND in Ethernet Virtual Private Networks*
RFC 9251, *Internet Group Management Protocol (IGMP) and Multicast Listener Discovery (MLD) Proxies for Ethernet VPN (EVPN)*

3.7 gRPC Remote Procedure Calls (gRPC)

cert.proto version 0.1.0, *gNOI Certificate Management Service*

file.proto version 0.1.0, *gNOI File Service*

gnmi.proto version 0.8.0, *gNMI Service Specification*

gnmi_ext.proto, *gNMI Commit Confirmed Extension*

gnmi_ext.proto, *gNMI Config Subscription Extension*

gnmi_ext.proto, *gNMI Depth Extension*

system.proto version 1.0.0, *gNOI System Service*

tunnel.proto version 0.2, *gRPC Tunnel Service*

PROTOCOL-HTTP2, *gRPC over HTTP2*

3.8 Intermediate System to Intermediate System (IS-IS)

draft-ietf-isis-mi-02, *IS-IS Multi-Instance*

draft-ietf-lsr-igp-ureach-prefix-announce-01, *IGP Unreachable Prefix Announcement* – without U-Flag and UP-Flag

draft-kaplan-isis-ext-eth-02, *Extended Ethernet Frame Size Support*

ISO/IEC 10589:2002 Second Edition, *Intermediate system to Intermediate system intra-domain routing information exchange protocol for use in conjunction with the protocol for providing the connectionless-mode Network Service (ISO 8473)*

RFC 1195, *Use of OSI IS-IS for Routing in TCP/IP and Dual Environments*

RFC 2973, *IS-IS Mesh Groups*

RFC 3359, *Reserved Type, Length and Value (TLV) Codepoints in Intermediate System to Intermediate System*

RFC 3719, *Recommendations for Interoperable Networks using Intermediate System to Intermediate System (IS-IS)*

RFC 3787, *Recommendations for Interoperable IP Networks using Intermediate System to Intermediate System (IS-IS)*

RFC 5120, *M-ISIS: Multi Topology (MT) Routing in IS-IS*

RFC 5130, *A Policy Control Mechanism in IS-IS Using Administrative Tags*

RFC 5301, *Dynamic Hostname Exchange Mechanism for IS-IS*

RFC 5302, *Domain-wide Prefix Distribution with Two-Level IS-IS*

RFC 5303, *Three-Way Handshake for IS-IS Point-to-Point Adjacencies*

RFC 5304, *IS-IS Cryptographic Authentication*

RFC 5305, *IS-IS Extensions for Traffic Engineering TE*

RFC 5306, *Restart Signaling for IS-IS – helper mode*

RFC 5308, *Routing IPv6 with IS-IS*
RFC 5309, *Point-to-Point Operation over LAN in Link State Routing Protocols*
RFC 5310, *IS-IS Generic Cryptographic Authentication*
RFC 5311, *Simplified Extension of Link State PDU (LSP) Space for IS-IS*
RFC 6213, *IS-IS BFD-Enabled TLV*
RFC 6232, *Purge Originator Identification TLV for IS-IS*
RFC 6233, *IS-IS Registry Extension for Purges*
RFC 7775, *IS-IS Route Preference for Extended IP and IPv6 Reachability*
RFC 7794, *IS-IS Prefix Attributes for Extended IPv4 and IPv6 Reachability* – sections 2.1 and 2.3
RFC 7981, *IS-IS Extensions for Advertising Router Information*
RFC 7987, *IS-IS Minimum Remaining Lifetime*
RFC 8202, *IS-IS Multi-Instance* – single topology
RFC 8570, *IS-IS Traffic Engineering (TE) Metric Extensions* – Min/Max Unidirectional Link Delay metric for flex-algo, RSVP, SR-TE
RFC 8919, *IS-IS Application-Specific Link Attributes*
RFC 9885, *Multi-Part TLVs in IS-IS*

3.9 Internet Protocol (IP) general

RFC 768, *User Datagram Protocol*
RFC 793, *Transmission Control Protocol*
RFC 854, *Telnet Protocol Specifications*
RFC 1350, *The TFTP Protocol (revision 2)*
RFC 2784, *Generic Routing Encapsulation (GRE)*
RFC 3164, *The BSD syslog Protocol*
RFC 4250, *The Secure Shell (SSH) Protocol Assigned Numbers*
RFC 4251, *The Secure Shell (SSH) Protocol Architecture*
RFC 4252, *The Secure Shell (SSH) Authentication Protocol* – publickey, password
RFC 4253, *The Secure Shell (SSH) Transport Layer Protocol*
RFC 4254, *The Secure Shell (SSH) Connection Protocol*
RFC 4511, *Lightweight Directory Access Protocol (LDAP): The Protocol*
RFC 4513, *Lightweight Directory Access Protocol (LDAP): Authentication Methods and Security Mechanisms* – TLS
RFC 4632, *Classless Inter-domain Routing (CIDR): The Internet Address Assignment and Aggregation Plan*
RFC 5082, *The Generalized TTL Security Mechanism (GTSM)*
RFC 5246, *The Transport Layer Security (TLS) Protocol Version 1.2* – TLS client, RSA public key

RFC 5289, *TLS Elliptic Curve Cipher Suites with SHA-256/384 and AES Galois Counter Mode (GCM)*
RFC 5425, *Transport Layer Security (TLS) Transport Mapping for Syslog – RFC 3164 with TLS*
RFC 5656, *Elliptic Curve Algorithm Integration in the Secure Shell Transport Layer – ECDSA*
RFC 5925, *The TCP Authentication Option*
RFC 5926, *Cryptographic Algorithms for the TCP Authentication Option (TCP-AO)*
RFC 6398, *IP Router Alert Considerations and Usage – MLD*
RFC 6528, *Defending against Sequence Number Attacks*
RFC 8446, *The Transport Layer Security (TLS) Protocol Version 1.3*
RFC 8907, *The Terminal Access Controller Access-Control System Plus (TACACS+) Protocol*

3.10 Internet Protocol (IP) multicast

RFC 1112, *Host Extensions for IP Multicasting*
RFC 2236, *Internet Group Management Protocol, Version 2*
RFC 2365, *Administratively Scoped IP Multicast*
RFC 2375, *IPv6 Multicast Address Assignments*
RFC 2710, *Multicast Listener Discovery (MLD) for IPv6*
RFC 3376, *Internet Group Management Protocol, Version 3*
RFC 3446, *Anycast Rendezvous Point (RP) mechanism using Protocol Independent Multicast (PIM) and Multicast Source Discovery Protocol (MSDP)*
RFC 3590, *Source Address Selection for the Multicast Listener Discovery (MLD) Protocol*
RFC 3618, *Multicast Source Discovery Protocol (MSDP)*
RFC 3810, *Multicast Listener Discovery Version 2 (MLDv2) for IPv6*
RFC 4541, *Considerations for Internet Group Management Protocol (IGMP) and Multicast Listener Discovery (MLD) Snooping Switches*
RFC 4604, *Using Internet Group Management Protocol Version 3 (IGMPv3) and Multicast Listener Discovery Protocol Version 2 (MLDv2) for Source-Specific Multicast*
RFC 4610, *Anycast-RP Using Protocol Independent Multicast (PIM)*
RFC 4611, *Multicast Source Discovery Protocol (MSDP) Deployment Scenarios*
RFC 5059, *Bootstrap Router (BSR) Mechanism for Protocol Independent Multicast (PIM)*
RFC 5186, *Internet Group Management Protocol Version 3 (IGMPv3) / Multicast Listener Discovery Version 2 (MLDv2) and Multicast Routing Protocol Interaction*
RFC 5384, *The Protocol Independent Multicast (PIM) Join Attribute Format*
RFC 7761, *Protocol Independent Multicast - Sparse Mode (PIM-SM): Protocol Specification (Revised)*
RFC 8487, *Mtrace Version 2: Traceroute Facility for IP Multicast*

3.11 Internet Protocol (IP) version 4

RFC 791, *Internet Protocol*
RFC 792, *Internet Control Message Protocol*
RFC 826, *An Ethernet Address Resolution Protocol*
RFC 1034, *Domain Names - Concepts and Facilities*
RFC 1035, *Domain Names - Implementation and Specification*
RFC 1191, *Path MTU Discovery* – router specification
RFC 1519, *Classless Inter-Domain Routing (CIDR): an Address Assignment and Aggregation Strategy*
RFC 1812, *Requirements for IPv4 Routers*
RFC 1918, *Address Allocation for Private Internets*
RFC 2131, *Dynamic Host Configuration Protocol*; Relay only
RFC 2132, *DHCP Options and BOOTP Vendor Extensions* – DHCP
RFC 2401, *Security Architecture for Internet Protocol*
RFC 3021, *Using 31-Bit Prefixes on IPv4 Point-to-Point Links*
RFC 3046, *DHCP Relay Agent Information Option (Option 82)*
RFC 3768, *Virtual Router Redundancy Protocol (VRRP)*
RFC 4884, *Extended ICMP to Support Multi-Part Messages* – ICMPv4 and ICMPv6 Time Exceeded

3.12 Internet Protocol (IP) version 6

RFC 2464, *Transmission of IPv6 Packets over Ethernet Networks*
RFC 2529, *Transmission of IPv6 over IPv4 Domains without Explicit Tunnels*
RFC 3122, *Extensions to IPv6 Neighbor Discovery for Inverse Discovery Specification*
RFC 3315, *Dynamic Host Configuration Protocol for IPv6 (DHCPv6)*
RFC 3587, *IPv6 Global Unicast Address Format*
RFC 3596, *DNS Extensions to Support IP version 6*
RFC 3633, *IPv6 Prefix Options for Dynamic Host Configuration Protocol (DHCP) version 6*
RFC 3736, *Stateless Dynamic Host Configuration Protocol (DHCP) Service for IPv6*
RFC 3971, *SEcure Neighbor Discovery (SEND)*
RFC 4007, *IPv6 Scoped Address Architecture*
RFC 4191, *Default Router Preferences and More-Specific Routes* – Default Router Preference
RFC 4193, *Unique Local IPv6 Unicast Addresses*
RFC 4291, *Internet Protocol Version 6 (IPv6) Addressing Architecture*
RFC 4443, *Internet Control Message Protocol (ICMPv6) for the Internet Protocol Version 6 (IPv6) Specification*

RFC 4861, *Neighbor Discovery for IP version 6 (IPv6)*
RFC 5095, *Deprecation of Type 0 Routing Headers in IPv6*
RFC 5722, *Handling of Overlapping IPv6 Fragments*
RFC 5798, *Virtual Router Redundancy Protocol (VRRP) Version 3 for IPv4 and IPv6 – IPv6*
RFC 5952, *A Recommendation for IPv6 Address Text Representation*
RFC 6164, *Using 127-Bit IPv6 Prefixes on Inter-Router Links*
RFC 8021, *Generation of IPv6 Atomic Fragments Considered Harmful*
RFC 8200, *Internet Protocol, Version 6 (IPv6) Specification*

3.13 Internet Protocol Security (IPsec)

draft-ietf-ipsec-isakmp-mode-cfg-05, *The ISAKMP Configuration Method*
draft-ietf-ipsec-isakmp-xauth-06, *Extended Authentication within ISAKMP/Oakley (XAUTH)*
RFC 2401, *Security Architecture for the Internet Protocol*
RFC 2403, *The Use of HMAC-MD5-96 within ESP and AH*
RFC 2404, *The Use of HMAC-SHA-1-96 within ESP and AH*
RFC 2405, *The ESP DES-CBC Cipher Algorithm With Explicit IV*
RFC 2406, *IP Encapsulating Security Payload (ESP)*
RFC 2407, *IPsec Domain of Interpretation for ISAKMP (IPsec DoI)*
RFC 2408, *Internet Security Association and Key Management Protocol (ISAKMP)*
RFC 2409, *The Internet Key Exchange (IKE)*
RFC 2410, *The NULL Encryption Algorithm and Its Use With IPsec*
RFC 2560, *X.509 Internet Public Key Infrastructure Online Certificate Status Protocol - OCSP*
RFC 3526, *More Modular Exponential (MODP) Diffie-Hellman group for Internet Key Exchange (IKE)*
RFC 3566, *The AES-XCBC-MAC-96 Algorithm and Its Use With IPsec*
RFC 3602, *The AES-CBC Cipher Algorithm and Its Use with IPsec*
RFC 3706, *A Traffic-Based Method of Detecting Dead Internet Key Exchange (IKE) Peers*
RFC 3947, *Negotiation of NAT-Traversal in the IKE*
RFC 3948, *UDP Encapsulation of IPsec ESP Packets*
RFC 4106, *The Use of Galois/Counter Mode (GCM) in IPsec ESP*
RFC 4109, *Algorithms for Internet Key Exchange version 1 (IKEv1)*
RFC 4301, *Security Architecture for the Internet Protocol*
RFC 4303, *IP Encapsulating Security Payload*
RFC 4307, *Cryptographic Algorithms for Use in the Internet Key Exchange Version 2 (IKEv2)*
RFC 4308, *Cryptographic Suites for IPsec*
RFC 4434, *The AES-XCBC-PRF-128 Algorithm for the Internet Key Exchange Protocol (IKE)*

RFC 4543, *The Use of Galois Message Authentication Code (GMAC) in IPsec ESP and AH*

RFC 4754, *IKE and IKEv2 Authentication Using the Elliptic Curve Digital Signature Algorithm (ECDSA)*

RFC 4835, *Cryptographic Algorithm Implementation Requirements for Encapsulating Security Payload (ESP) and Authentication Header (AH)*

RFC 4868, *Using HMAC-SHA-256, HMAC-SHA-384, and HMAC-SHA-512 with IPsec*

RFC 4945, *The Internet IP Security PKI Profile of IKEv1/ISAKMP, IKEv2 and PKIX*

RFC 5019, *The Lightweight Online Certificate Status Protocol (OCSP) Profile for High-Volume Environments*

RFC 5282, *Using Authenticated Encryption Algorithms with the Encrypted Payload of the IKEv2 Protocol*

RFC 5903, *ECP Groups for IKE and IKEv2*

RFC 5996, *Internet Key Exchange Protocol Version 2 (IKEv2)*

RFC 5998, *An Extension for EAP-Only Authentication in IKEv2*

RFC 6379, *Suite B Cryptographic Suites for IPsec*

RFC 6380, *Suite B Profile for Internet Protocol Security (IPsec)*

RFC 6960, *X.509 Internet Public Key Infrastructure Online Certificate Status Protocol - OCSP*

RFC 7296, *Internet Key Exchange Protocol Version 2 (IKEv2)*

RFC 7321, *Cryptographic Algorithm Implementation Requirements and Usage Guidance for Encapsulating Security Payload (ESP) and Authentication Header (AH)*

RFC 7383, *Internet Key Exchange Protocol Version 2 (IKEv2) Message Fragmentation*

RFC 7427, *Signature Authentication in the Internet Key Exchange Version 2 (IKEv2)*

RFC 8784, *Mixing Preshared Keys in the Internet Key Exchange Protocol Version 2 (IKEv2) for Post-quantum Security*

3.14 Label Distribution Protocol (LDP)

draft-pdutta-mpls-ldp-adj-capability-00, *LDP Adjacency Capabilities*

draft-pdutta-mpls-ldp-v2-00, *LDP Version 2*

RFC 3037, *LDP Applicability*

RFC 3478, *Graceful Restart Mechanism for Label Distribution Protocol – helper mode*

RFC 5036, *LDP Specification*

RFC 5283, *LDP Extension for Inter-Area Label Switched Paths (LSPs)*

RFC 5443, *LDP IGP Synchronization*

RFC 5561, *LDP Capabilities*

RFC 5919, *Signaling LDP Label Advertisement Completion*

3.15 Multiprotocol Label Switching (MPLS)

RFC 3031, *Multiprotocol Label Switching Architecture*

RFC 3032, *MPLS Label Stack Encoding*

RFC 3270, *Multi-Protocol Label Switching (MPLS) Support of Differentiated Services – E-LSP*

RFC 3443, *Time To Live (TTL) Processing in Multi-Protocol Label Switching (MPLS) Networks*

RFC 5332, *MPLS Multicast Encapsulations*

RFC 5884, *Bidirectional Forwarding Detection (BFD) for MPLS Label Switched Paths (LSPs)*

RFC 6424, *Mechanism for Performing Label Switched Path Ping (LSP Ping) over MPLS Tunnels*

RFC 7308, *Extended Administrative Groups in MPLS Traffic Engineering (MPLS-TE)*

RFC 7746, *Label Switched Path (LSP) Self-Ping*

3.16 Network Address Translation (NAT)

RFC 4787, *Network Address Translation (NAT) Behavioral Requirements for Unicast UDP*

RFC 5382, *NAT Behavioral Requirements for TCP*

RFC 5508, *NAT Behavioral Requirements for ICMP*

3.17 Network Configuration Protocol (NETCONF)

RFC 5277, *NETCONF Event Notifications*

RFC 6020, *YANG - A Data Modeling Language for the Network Configuration Protocol (NETCONF)*

RFC 6022, *YANG Module for NETCONF Monitoring*

RFC 6241, *Network Configuration Protocol (NETCONF)*

RFC 6242, *Using the NETCONF Protocol over Secure Shell (SSH)*

RFC 6243, *With-defaults Capability for NETCONF*

RFC 8071, *NETCONF Call Home and RESTCONF Call Home – NETCONF*

RFC 8342, *Network Management Datastore Architecture (NMDA) – Startup, Candidate, Running and Intended datastores*

RFC 8525, *YANG Library*

RFC 8526, *NETCONF Extensions to Support the Network Management Datastore Architecture – <get-data> operation*

3.18 Open Shortest Path First (OSPF)

RFC 1765, *OSPF Database Overflow*

RFC 2328, *OSPF Version 2*
RFC 3101, *The OSPF Not-So-Stubby Area (NSSA) Option*
RFC 3509, *Alternative Implementations of OSPF Area Border Routers*
RFC 3623, *Graceful OSPF Restart Graceful OSPF Restart – helper mode*
RFC 3630, *Traffic Engineering (TE) Extensions to OSPF Version 2*
RFC 4811, *OSPF Out-of-Band Link State Database (LSDB) Resynchronization – OSPFv2*
RFC 4812, *OSPF Restart Signaling – OSPFv2*
RFC 4222, *Prioritized Treatment of Specific OSPF Version 2 Packets and Congestion Avoidance*
RFC 4552, *Authentication/Confidentiality for OSPFv3*
RFC 4576, *Using a Link State Advertisement (LSA) Options Bit to Prevent Looping in BGP/MPLS IP Virtual Private Networks (VPNs)*
RFC 4577, *OSPF as the Provider/Customer Edge Protocol for BGP/MPLS IP Virtual Private Networks (VPNs)*
RFC 5185, *OSPF Multi-Area Adjacency*
RFC 5243, *OSPF Database Exchange Summary List Optimization*
RFC 5250, *The OSPF Opaque LSA Option*
RFC 5309, *Point-to-Point Operation over LAN in Link State Routing Protocols*
RFC 5340, *OSPF for IPv6*
RFC 5642, *Dynamic Hostname Exchange Mechanism for OSPF*
RFC 5709, *OSPFv2 HMAC-SHA Cryptographic Authentication*
RFC 5838, *Support of Address Families in OSPFv3*
RFC 6549, *OSPFv2 Multi-Instance Extensions*
RFC 6987, *OSPF Stub Router Advertisement*
RFC 7471, *OSPF Traffic Engineering (TE) Metric Extensions – Min/Max Unidirectional Link Delay metric for flex-algo, RSVP, SR-TE*
RFC 7684, *OSPFv2 Prefix/Link Attribute Advertisement*
RFC 7770, *Extensions to OSPF for Advertising Optional Router Capabilities*
RFC 8362, *OSPFv3 Link State Advertisement (LSA) Extensibility*
RFC 8920, *OSPF Application-Specific Link Attributes*

3.19 Path Computation Element Protocol (PCEP)

draft-alvarez-pce-path-profiles-04, *PCE Path Profiles*
draft-ietf-pce-binding-label-sid-15, *Carrying Binding Label/Segment Identifier (SID) in PCE-based Networks. – MPLS binding SIDs*
RFC 5440, *Path Computation Element (PCE) Communication Protocol (PCEP)*
RFC 8231, *Path Computation Element Communication Protocol (PCEP) Extensions for Stateful PCE*

RFC 8253, *PCEPS: Usage of TLS to Provide a Secure Transport for the Path Computation Element Communication Protocol (PCEP)*

RFC 8281, *PCEP Extensions for PCE-initiated LSP Setup in a Stateful PCE Model*

RFC 8408, *Conveying Path Setup Type in PCE Communication Protocol (PCEP) Messages*

RFC 8664, *Path Computation Element Communication Protocol (PCEP) Extensions for Segment Routing*

3.20 Pseudowire (PW)

draft-ietf-l2vpn-vpws-iw-oam-04, *OAM Procedures for VPWS Interworking*

MFA Forum 12.0.0, *Multiservice Interworking - Ethernet over MPLS*

MFA Forum 13.0.0, *Fault Management for Multiservice Interworking v1.0*

MFA Forum 16.0.0, *Multiservice Interworking - IP over MPLS*

RFC 3916, *Requirements for Pseudo-Wire Emulation Edge-to-Edge (PWE3)*

RFC 3985, *Pseudo Wire Emulation Edge-to-Edge (PWE3)*

RFC 4385, *Pseudo Wire Emulation Edge-to-Edge (PWE3) Control Word for Use over an MPLS PSN*

RFC 4446, *IANA Allocations for Pseudowire Edge to Edge Emulation (PWE3)*

RFC 4447, *Pseudowire Setup and Maintenance Using the Label Distribution Protocol (LDP)*

RFC 4448, *Encapsulation Methods for Transport of Ethernet over MPLS Networks*

RFC 5085, *Pseudowire Virtual Circuit Connectivity Verification (VCCV): A Control Channel for Pseudowires*

RFC 5659, *An Architecture for Multi-Segment Pseudowire Emulation Edge-to-Edge*

RFC 5885, *Bidirectional Forwarding Detection (BFD) for the Pseudowire Virtual Circuit Connectivity Verification (VCCV)*

RFC 6073, *Segmented Pseudowire*

RFC 6310, *Pseudowire (PW) Operations, Administration, and Maintenance (OAM) Message Mapping*

RFC 6391, *Flow-Aware Transport of Pseudowires over an MPLS Packet Switched Network*

RFC 6575, *Address Resolution Protocol (ARP) Mediation for IP Interworking of Layer 2 VPNs*

RFC 6718, *Pseudowire Redundancy*

RFC 6829, *Label Switched Path (LSP) Ping for Pseudowire Forwarding Equivalence Classes (FECs) Advertised over IPv6*

RFC 6870, *Pseudowire Preferential Forwarding Status bit*

RFC 7023, *MPLS and Ethernet Operations, Administration, and Maintenance (OAM) Interworking*

RFC 7267, *Dynamic Placement of Multi-Segment Pseudowires*

RFC 7392, *Explicit Path Routing for Dynamic Multi-Segment Pseudowires – ER-TLV and ER-HOP IPv4 Prefix*

RFC 8395, *Extensions to BGP-Signaled Pseudowires to Support Flow-Aware Transport Labels*

3.21 Quality of Service (QoS)

RFC 2430, *A Provider Architecture for Differentiated Services and Traffic Engineering (PASTE)*
RFC 2474, *Definition of the Differentiated Services Field (DS Field) in the IPv4 and IPv6 Headers*
RFC 2597, *Assured Forwarding PHB Group*
RFC 3140, *Per Hop Behavior Identification Codes*
RFC 3246, *An Expedited Forwarding PHB (Per-Hop Behavior)*

3.22 Remote Authentication Dial In User Service (RADIUS)

draft-oscca-cfrg-sm3-02, *The SM3 Cryptographic Hash Function*
RFC 2865, *Remote Authentication Dial In User Service (RADIUS)*
RFC 2866, *RADIUS Accounting*
RFC 3162, *RADIUS and IPv6*
RFC 6613, *RADIUS over TCP – with TLS*
RFC 6614, *Transport Layer Security (TLS) Encryption for RADIUS*
RFC 6929, *Remote Authentication Dial-In User Service (RADIUS) Protocol Extensions*

3.23 Resource Reservation Protocol - Traffic Engineering (RSVP-TE)

RFC 2702, *Requirements for Traffic Engineering over MPLS*
RFC 2747, *RSVP Cryptographic Authentication*
RFC 2961, *RSVP Refresh Overhead Reduction Extensions*
RFC 3097, *RSVP Cryptographic Authentication -- Updated Message Type Value*
RFC 3209, *RSVP-TE: Extensions to RSVP for LSP Tunnels*
RFC 4124, *Protocol Extensions for Support of Diffserv-aware MPLS Traffic Engineering*
RFC 4125, *Maximum Allocation Bandwidth Constraints Model for Diffserv-aware MPLS Traffic Engineering*
RFC 4127, *Russian Dolls Bandwidth Constraints Model for Diffserv-aware MPLS Traffic Engineering*
RFC 4561, *Definition of a Record Route Object (RRO) Node-Id Sub-Object*
RFC 5817, *Graceful Shutdown in MPLS and Generalized MPLS Traffic Engineering Networks*

3.24 Routing Information Protocol (RIP)

RFC 1058, *Routing Information Protocol*
RFC 2080, *RIPng for IPv6*
RFC 2082, *RIP-2 MD5 Authentication*

RFC 2453, *RIP Version 2*

3.25 Segment Routing (SR)

RFC 8287, *Label Switched Path (LSP) Ping/Traceroute for Segment Routing (SR) IGP-Prefix and IGP-Adjacency Segment Identifiers (SIDs) with MPLS Data Planes*

RFC 8426, *Recommendations for RSVP-TE and Segment Routing (SR) Label Switched Path (LSP) Coexistence*

RFC 8476, *Signaling Maximum SID Depth (MSD) Using OSPF – node MSD*

RFC 8491, *Signaling Maximum SID Depth (MSD) Using IS-IS – node MSD*

RFC 8660, *Segment Routing with the MPLS Data Plane*

RFC 8661, *Segment Routing MPLS Interworking with LDP*

RFC 8665, *OSPF Extensions for Segment Routing*

RFC 8667, *IS-IS Extensions for Segment Routing*

RFC 8669, *Segment Routing Prefix Segment Identifier Extensions for BGP*

RFC 9256, *Segment Routing Policy Architecture*

RFC 9350, *IGP Flexible Algorithm*

3.26 Simple Network Management Protocol (SNMP)

draft-blumenthal-aes-usm-04, *The AES Cipher Algorithm in the SNMP's User-based Security Model – CFB128-AES-192 and CFB128-AES-256*

draft-ietf-isis-wg-mib-06, *Management Information Base for Intermediate System to Intermediate System (IS-IS)*

draft-ietf-mpls-ldp-mib-07, *Definitions of Managed Objects for the Multiprotocol Label Switching, Label Distribution Protocol (LDP)*

draft-ietf-mpls-lsr-mib-06, *Multiprotocol Label Switching (MPLS) Label Switching Router (LSR) Management Information Base Using SMIv2*

draft-ietf-mpls-te-mib-04, *Multiprotocol Label Switching (MPLS) Traffic Engineering Management Information Base*

draft-ietf-ospf-mib-update-08, *OSPF Version 2 Management Information Base*

draft-ietf-vrrp-unified-mib-06, *Definitions of Managed Objects for the VRRP over IPv4 and IPv6 – IPv6*

ESO-CONSORTIUM-MIB revision 200406230000Z, *esoConsortiumMIB*

IANA-ADDRESS-FAMILY-NUMBERS-MIB revision 200203140000Z, *ianaAddressFamilyNumbers*

IANAifType-MIB revision 200505270000Z, *ianaifType*

IANA-RTPROTO-MIB revision 200009260000Z, *ianaRtProtoMIB*

IEEE8021-CFM-MIB revision 200706100000Z, *ieee8021CfmMib*

IEEE8021-PAE-MIB revision 200101160000Z, *ieee8021paeMIB*

IEEE8023-LAG-MIB revision 200006270000Z, *lagMIB*
LLDP-MIB revision 200505060000Z, *lldpMIB*
RFC 1157, *A Simple Network Management Protocol (SNMP)*
RFC 1212, *Concise MIB Definitions*
RFC 1215, *A Convention for Defining Traps for use with the SNMP*
RFC 1724, *RIP Version 2 MIB Extension*
RFC 1901, *Introduction to Community-based SNMPv2*
RFC 2021, *Remote Network Monitoring Management Information Base Version 2 using SMIv2*
RFC 2206, *RSVP Management Information Base using SMIv2*
RFC 2578, *Structure of Management Information Version 2 (SMIv2)*
RFC 2579, *Textual Conventions for SMIv2*
RFC 2580, *Conformance Statements for SMIv2*
RFC 2787, *Definitions of Managed Objects for the Virtual Router Redundancy Protocol*
RFC 2819, *Remote Network Monitoring Management Information Base*
RFC 2856, *Textual Conventions for Additional High Capacity Data Types*
RFC 2863, *The Interfaces Group MIB*
RFC 2864, *The Inverted Stack Table Extension to the Interfaces Group MIB*
RFC 2933, *Internet Group Management Protocol MIB*
RFC 3014, *Notification Log MIB*
RFC 3273, *Remote Network Monitoring Management Information Base for High Capacity Networks*
RFC 3410, *Introduction and Applicability Statements for Internet Standard Management Framework*
RFC 3430, *Simple Network Management Protocol (SNMP) over Transmission Control Protocol (TCP) Transport Mapping*
RFC 3411, *An Architecture for Describing Simple Network Management Protocol (SNMP) Management Frameworks*
RFC 3412, *Message Processing and Dispatching for the Simple Network Management Protocol (SNMP)*
RFC 3413, *Simple Network Management Protocol (SNMP) Applications*
RFC 3414, *User-based Security Model (USM) for version 3 of the Simple Network Management Protocol (SNMPv3)*
RFC 3415, *View-based Access Control Model (VACM) for the Simple Network Management Protocol (SNMP)*
RFC 3416, *Version 2 of the Protocol Operations for the Simple Network Management Protocol (SNMP)*
RFC 3417, *Transport Mappings for the Simple Network Management Protocol (SNMP) – SNMP over UDP over IPv4*
RFC 3418, *Management Information Base (MIB) for the Simple Network Management Protocol (SNMP)*
RFC 3419, *Textual Conventions for Transport Addresses*
RFC 3434, *Remote Monitoring MIB Extensions for High Capacity Alarms*

RFC 3584, *Coexistence between Version 1, Version 2, and Version 3 of the Internet-standard Network Management Framework*

RFC 3593, *Textual Conventions for MIB Modules Using Performance History Based on 15 Minute Intervals*

RFC 3635, *Definitions of Managed Objects for the Ethernet-like Interface Types*

RFC 3826, *The Advanced Encryption Standard (AES) Cipher Algorithm in the SNMP User-based Security Model*

RFC 3877, *Alarm Management Information Base (MIB)*

RFC 4001, *Textual Conventions for Internet Network Addresses*

RFC 4022, *Management Information Base for the Transmission Control Protocol (TCP)*

RFC 4113, *Management Information Base for the User Datagram Protocol (UDP)*

RFC 4273, *Definitions of Managed Objects for BGP-4*

RFC 4292, *IP Forwarding Table MIB*

RFC 4293, *Management Information Base for the Internet Protocol (IP)*

RFC 4878, *Definitions and Managed Objects for Operations, Administration, and Maintenance (OAM) Functions on Ethernet-Like Interfaces*

RFC 7420, *Path Computation Element Communication Protocol (PCEP) Management Information Base (MIB) Module*

RFC 7630, *HMAC-SHA-2 Authentication Protocols in the User-based Security Model (USM) for SNMPv3*

3.27 Storage sanitization

IEEE 2883-2022, *IEEE Standard for Sanitizing Storage* – clear sanitization method for CF, MMC, SSD, SD, USB

3.28 Timing

RFC 3339, *Date and Time on the Internet: Timestamps*

RFC 5905, *Network Time Protocol Version 4: Protocol and Algorithms Specification*

RFC 8573, *Message Authentication Code for the Network Time Protocol*

3.29 Two-Way Active Measurement Protocol (TWAMP)

RFC 5357, *A Two-Way Active Measurement Protocol (TWAMP) – server, unauthenticated mode*

RFC 5938, *Individual Session Control Feature for the Two-Way Active Measurement Protocol (TWAMP)*

RFC 6038, *Two-Way Active Measurement Protocol (TWAMP) Reflect Octets and Symmetrical Size Features*

RFC 8545, *Well-Known Port Assignments for the One-Way Active Measurement Protocol (OWAMP) and the Two-Way Active Measurement Protocol (TWAMP) – TWAMP*

RFC 8762, *Simple Two-Way Active Measurement Protocol* – unauthenticated

RFC 8972, *Simple Two-Way Active Measurement Protocol Optional Extensions* – unauthenticated

RFC 9503, *Simple Two-Way Active Measurement Protocol (STAMP) Extensions for Segment Routing Networks* – excluding Sections 3, 4.1.2 and 4.1.3

3.30 Virtual Private LAN Service (VPLS)

RFC 4761, *Virtual Private LAN Service (VPLS) Using BGP for Auto-Discovery and Signaling*

RFC 4762, *Virtual Private LAN Service (VPLS) Using Label Distribution Protocol (LDP) Signaling*

RFC 5501, *Requirements for Multicast Support in Virtual Private LAN Services*

RFC 6074, *Provisioning, Auto-Discovery, and Signaling in Layer 2 Virtual Private Networks (L2VPNs)*

RFC 7117, *Multicast in Virtual Private LAN Service (VPLS)*

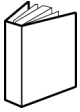
3.31 Yet Another Next Generation (YANG)

RFC 6991, *Common YANG Data Types*

RFC 7950, *The YANG 1.1 Data Modeling Language*

RFC 7951, *JSON Encoding of Data Modeled with YANG*

Customer document and product support



Customer documentation

[Customer documentation welcome page](#)



Technical support

[Product support portal](#)



Documentation feedback

[Customer documentation feedback](#)