



7705 Service Aggregation Router Gen 2

Release 26.7.R1

Layer 3 Services Guide: IES and VPRN

3HE 29559 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.....	10
List of figures.....	11
1 Getting started.....	14
1.1 About this guide.....	14
1.2 Platforms and terminology.....	14
1.3 Conventions.....	15
1.3.1 Precautionary and information messages.....	15
1.3.2 Options or substeps in procedures and sequential workflows.....	15
2 Internet Enhanced Service.....	17
2.1 IES overview.....	17
2.2 IES features.....	17
2.2.1 IP interfaces.....	18
2.2.1.1 Object grouping and state monitoring.....	18
2.2.2 SAPs.....	20
2.2.2.1 Encapsulations.....	20
2.2.2.2 Encapsulation.....	20
2.2.2.3 Shaping and bandwidth control.....	21
2.2.2.4 Lag considerations.....	22
2.2.2.5 Last mile packet size adjustment.....	22
2.2.2.6 Redundancy with pseudowire SAPs.....	23
2.2.2.7 Operational group support for PW ports.....	25
2.2.3 Routing protocols.....	27
2.2.3.1 CPE connectivity check.....	27
2.2.4 QoS policies.....	27
2.2.5 Filter policies.....	28
2.2.6 MPLS EL and hash label.....	28
2.2.7 Spoke SDPs.....	28
2.2.8 DHCP.....	30
2.2.8.1 DHCP client for IES.....	30
2.2.8.2 DHCP relay for IES.....	30
2.2.8.3 Local DHCP server for IES.....	30

2.3	Configuring an IES service with CLI.....	30
2.3.1	Basic configuration.....	31
2.3.2	Common configuration tasks.....	31
2.3.3	Configuring IES components.....	32
2.3.3.1	Configuring an IES service.....	32
2.3.3.2	Configuring IES subscriber interfaces.....	33
2.3.3.3	Configuring IES interfaces.....	34
2.3.3.4	Configuring a spoke SDP.....	34
2.3.3.5	Configuring a SAP.....	35
2.3.3.6	Configuring VRRP.....	36
2.3.3.7	Configuring IPsec.....	37
2.3.3.8	IGMP host tracking.....	38
2.4	Service management tasks.....	39
2.4.1	Modifying IES service command options.....	39
2.4.2	Deleting a spoke SDP.....	40
2.4.3	Deleting an IES service.....	40
2.4.4	Disabling an IES service.....	41
2.4.5	Re-enabling an IES service.....	41
3	VPRN service.....	43
3.1	VPRN service overview.....	43
3.1.1	Routing prerequisites.....	44
3.1.2	Core MP-BGP support.....	44
3.1.3	Route distinguishers.....	45
3.1.3.1	eiBGP load balancing.....	45
3.1.4	Route reflector.....	46
3.1.5	CE to PE route exchange.....	47
3.1.5.1	Route redistribution.....	47
3.1.5.2	CPE connectivity check.....	47
3.1.6	Constrained route distribution.....	50
3.1.6.1	Constrained VPN route distribution based on route targets.....	50
3.1.6.2	Configuring the route target address family.....	50
3.1.6.3	Originating RTC routes.....	50
3.1.6.4	Receiving and re-advertising RTC routes.....	51
3.1.6.5	Using RTC routes.....	52
3.1.7	BGP fast reroute in a VPRN.....	53

3.1.8	Export of inactive VPRN BGP routes.....	54
3.2	VPRN features.....	55
3.2.1	IP interfaces.....	56
3.2.1.1	Displaying QoS information associated with routes.....	56
3.2.1.2	Object grouping and state monitoring.....	57
3.2.1.3	VPRN IP interface applicability.....	57
3.2.2	SAPs.....	59
3.2.2.1	SAP encapsulations.....	59
3.2.2.2	TCP MSS adjustment filter on VPRN SAP interfaces.....	59
3.2.3	QoS policies.....	65
3.2.4	Filter policies.....	65
3.2.5	DSCP marking.....	65
3.2.5.1	Default DSCP mapping table.....	66
3.2.6	Configuration of TTL propagation for VPRN routes.....	67
3.2.7	CE to PE routing protocols.....	68
3.2.7.1	PE to PE tunneling mechanisms.....	68
3.2.7.2	Per VRF route limiting.....	69
3.2.8	Spoke SDPs.....	69
3.2.8.1	T-LDP status signaling for spoke-SDPs terminating on IES/VPRN.....	70
3.2.8.2	Spoke SDP redundancy into IES/VPRN.....	70
3.2.8.3	Weighted ECMP for spoke-SDPs terminating on IES/VPRN and R-VPLS interfaces.....	71
3.2.9	IP-VPNs.....	72
3.2.9.1	Using OSPF in IP-VPNs.....	72
3.2.10	IPCP subnet negotiation.....	72
3.2.11	Inter-AS VPRNs.....	73
3.2.11.1	Inter-AS option-A.....	73
3.2.11.2	Inter-AS Option B.....	73
3.2.11.3	Inter-AS option-C.....	76
3.2.12	VPRN label security at inter-AS boundary.....	77
3.2.12.1	Feature configuration.....	77
3.2.12.2	CPM behavior.....	78
3.2.12.3	Datapath forwarding behavior.....	79
3.2.13	CSC.....	80
3.2.13.1	Terminology.....	81
3.2.13.2	CSC connectivity models.....	81

3.2.13.3	CSC-PE configuration and operation.....	82
3.2.13.4	CSC interface.....	82
3.2.13.5	QoS.....	83
3.2.13.6	MPLS.....	85
3.2.13.7	CSC VPRN service configuration.....	85
3.2.14	Node management using VPRN.....	85
3.2.14.1	VPRN management.....	86
3.2.14.2	AAA management.....	87
3.2.14.3	SNMP management.....	88
3.2.14.4	Events and notifications.....	89
3.2.14.5	DNS resolution.....	90
3.2.15	Traffic leaking to GRT.....	90
3.2.15.1	Management via VPRN using GRT leaking.....	91
3.2.16	Traffic leaking from VPRN to GRT for IPv6.....	91
3.2.17	RIP metric propagation in VPRNs.....	92
3.2.18	NTP within a VPRN service.....	92
3.2.19	PTP within a VPRN service.....	93
3.2.20	VPN route label allocation.....	93
3.2.20.1	Configuring the service label mode.....	94
3.2.20.2	Restrictions and usage notes.....	94
3.2.21	VPRN Support for BGP FlowSpec.....	95
3.2.22	MPLS EL and hash label.....	95
3.2.23	LSP tagging for BGP next hops or prefixes and BGP-LU.....	96
3.2.24	Route leaking from GRT to VPRN instances.....	96
3.2.25	Class-based forwarding of VPN-v4/v6 prefixes over RSVP-TE or SR-TE LSPs.....	97
3.2.25.1	Feature configuration.....	97
3.2.25.2	Feature behavior.....	98
3.2.26	IP VPN independent domains using BGP attribute set.....	99
3.2.27	DHCP.....	101
3.2.27.1	DHCP client for VPRN.....	101
3.2.27.2	DHCP relay for VPRN.....	101
3.2.27.3	Local DHCP server for VPRN.....	101
3.3	QoS on ingress bindings.....	102
3.4	Multicast in IP-VPN applications.....	103
3.4.1	Use of data MDTs.....	105
3.4.2	Multicast protocols supported in the provider network.....	105

3.4.3	MVPN membership autodiscovery using BGP.....	106
3.4.4	PE-PE transmission of C-multicast routing using BGP.....	108
3.4.5	VRF route import extended community.....	108
3.4.6	Provider tunnel support.....	109
3.4.6.1	Point-to-Multipoint Inclusive (I-PMSI) and Selective (S-PMSI) Provider Multicast Service Interface.....	109
3.4.6.2	P2MP RSVP-TE I-PMSI and S-PMSI.....	109
3.4.6.3	P2MP LDP I-PMSI and S-PMSI.....	109
3.4.6.4	Wildcard (C-*, C-*) P2MP LSP S-PMSI.....	109
3.4.6.5	P2MP LSP S-PMSI.....	112
3.4.6.6	Dynamic multicast signaling over P2MP LDP in VRF.....	112
3.4.6.7	MVPN sender-only/receiver-only.....	114
3.4.6.8	S-PMSI trigger thresholds.....	115
3.4.6.9	Migration from existing Rosen implementation.....	116
3.4.6.10	Policy-based S-PMSI.....	116
3.4.6.11	Policy-based data MDT.....	119
3.4.7	MVPN (NG-MVPN) upstream multicast hop fast failover.....	120
3.4.8	Multicast VPN extranet.....	121
3.4.8.1	Multicast extranet for Rosen MVPN for PIM SSM.....	121
3.4.8.2	Multicast extranet for NG-MVPN for PIM SSM.....	122
3.4.8.3	Multicast extranet with per-group mapping for PIM SSM.....	123
3.4.8.4	Multicast GRT-source/VRF-receiver extranet with per group mapping for PIM SSM.....	125
3.4.8.5	Multicast extranet with per-group mapping for PIM ASM.....	126
3.4.9	Non-congruent unicast and multicast topologies for multicast VPN.....	127
3.4.10	Automatic discovery of Group-to-RP mappings (auto-RP).....	128
3.4.11	IPv6 MVPN support.....	128
3.4.12	Multicast core diversity for Rosen MDT_SAFI MVPNs.....	130
3.4.13	NG-MVPN core diversity.....	131
3.4.13.1	NG-MVPN to loopback interface.....	132
3.4.13.2	NG-MVPN core diversity.....	133
3.4.13.3	P2MP RSVP-TE core diversity with UFD for UMH redundancy.....	135
3.4.13.4	UFD packet generation.....	136
3.4.13.5	Configuration example.....	136
3.4.14	NG-MVPN multicast source geo-redundancy.....	142
3.4.15	Multicast core diversity for Rosen MDT SAFI MVPNs.....	144
3.4.16	Inter-AS MVPN.....	146

3.4.16.1	BGP connector attribute.....	146
3.4.16.2	PIM RPF vector.....	147
3.4.16.3	Inter-AS MVPN Option B.....	147
3.4.16.4	Inter-AS MVPN Option C.....	148
3.4.16.5	NG-MVPN non-segmented inter-AS solution.....	149
3.4.17	mLDP non-Segmented intra-AS (inter-area) MVPN solution.....	159
3.4.17.1	Intra-AS and Inter-AS Option B.....	159
3.4.17.2	MVPN next hop self on ABRs.....	159
3.4.18	UMH redundancy for bandwidth monitoring using a single IOM.....	164
3.4.18.1	Fault recovery mitigation at PMSI switchover time.....	165
3.4.18.2	S-PMSI behavior.....	165
3.4.18.3	Bandwidth monitoring on single IOMs.....	165
3.4.18.4	ASM behavior.....	166
3.4.18.5	Low traffic rate.....	166
3.4.18.6	Revertive timer.....	166
3.4.18.7	MVPN upstream PE fast failover.....	166
3.4.18.8	Multicast-only Fast Reroute.....	167
3.4.19	UMH redundancy for bandwidth monitoring using multiple IOMs.....	167
3.4.19.1	Implementation overview.....	168
3.4.19.2	Configuring bandwidth monitoring without redundant IOMs.....	172
3.4.19.3	Configuring redundant monitoring using a LAG.....	175
3.4.19.4	Viewing the port used for monitoring.....	177
3.5	FIB prioritization.....	178
3.6	Configuring a VPRN service using CLI.....	178
3.6.1	Basic configuration.....	178
3.6.2	Common configuration tasks.....	182
3.6.3	Configuring VPRN components.....	182
3.6.3.1	Creating a VPRN service.....	182
3.6.3.2	Configuring a global VPRN service.....	183
3.6.3.3	Configuring a VPRN log.....	184
3.6.3.4	Configuring VPRN protocols - PIM.....	186
3.7	Service management tasks.....	203
3.7.1	Modifying a VPRN service.....	203
3.7.2	Deleting a VPRN service.....	205
3.7.3	Disabling a VPRN service.....	206
3.7.4	Re-enabling a VPRN service.....	206

4	Standards and protocol support.....	208
4.1	Bidirectional Forwarding Detection (BFD).....	208
4.2	Border Gateway Protocol (BGP).....	208
4.3	Bridging and management.....	209
4.4	Certificate management.....	210
4.5	Ethernet.....	210
4.6	Ethernet VPN (EVPN).....	210
4.7	gRPC Remote Procedure Calls (gRPC).....	211
4.8	Intermediate System to Intermediate System (IS-IS).....	211
4.9	Internet Protocol (IP) general.....	212
4.10	Internet Protocol (IP) multicast.....	213
4.11	Internet Protocol (IP) version 4.....	214
4.12	Internet Protocol (IP) version 6.....	214
4.13	Internet Protocol Security (IPsec).....	215
4.14	Label Distribution Protocol (LDP).....	216
4.15	Multiprotocol Label Switching (MPLS).....	217
4.16	Network Address Translation (NAT).....	217
4.17	Network Configuration Protocol (NETCONF).....	217
4.18	Open Shortest Path First (OSPF).....	217
4.19	Path Computation Element Protocol (PCEP).....	218
4.20	Pseudowire (PW).....	219
4.21	Quality of Service (QoS).....	220
4.22	Remote Authentication Dial In User Service (RADIUS).....	220
4.23	Resource Reservation Protocol - Traffic Engineering (RSVP-TE).....	220
4.24	Routing Information Protocol (RIP).....	220
4.25	Segment Routing (SR).....	221
4.26	Simple Network Management Protocol (SNMP).....	221
4.27	Storage sanitization.....	223
4.28	Timing.....	223
4.29	Two-Way Active Measurement Protocol (TWAMP).....	223
4.30	Virtual Private LAN Service (VPLS).....	224
4.31	Yet Another Next Generation (YANG).....	224

List of tables

Table 1: Platforms and terminology.....14

Table 2: Packet sizes used for pseudowire SAPs..... 23

Table 3: BGP fast reroute scenarios (VPRN context).....54

Table 4: DSCP/FC marking..... 65

Table 5: Service labels distributed in service label per next hop mode.....94

Table 6: Supported configuration combinations..... 106

Table 7: Supported configuration combinations..... 107

Table 8: Recursive opaque types..... 151

List of figures

Figure 1: Internet Enhanced Service.....	17
Figure 2: PW SAP egress scheduling hierarchy options.....	22
Figure 3: Dual homing into multiple Layer 3 PEs.....	23
Figure 4: Master-standby PW redundancy.....	24
Figure 5: SDP-ID and VC label service identifiers.....	29
Figure 6: IES spoke-SDP termination.....	29
Figure 7: Virtual Private Routed Network.....	44
Figure 8: Route distinguisher.....	45
Figure 9: Basic eiBGP topology.....	46
Figure 10: Extranet load balancing.....	46
Figure 11: Directly connected IP target (MD-CLI).....	48
Figure 12: Directly connected IP target (classic CLI).....	48
Figure 13: Multiple hops to IP target (MD-CLI).....	49
Figure 14: Multiple hops to IP target (classic CLI).....	49
Figure 15: SDP-ID and VC label service identifiers.....	69
Figure 16: Active/standby VRF using resilient Layer 2 circuits.....	70
Figure 17: Spoke SDP redundancy model.....	71
Figure 18: CPEs network up-link mode.....	72
Figure 19: Inter-AS option-A: VRF-to-VRF model.....	73
Figure 20: Identical system IP on multiple PEs (Option B).....	74
Figure 21: Non-segmented mLDP PMSI establishment (Option B).....	75

Figure 22: Non-segmented mLDP C-multicast exchange (Option B).....	76
Figure 23: Inter-AS option-C.....	77
Figure 24: Carrier Supporting Carrier reference diagram.....	81
Figure 25: VRF network example.....	86
Figure 26: RIP metric propagation in VPRNs.....	92
Figure 27: Ingress QoS control on VPRN bindings.....	102
Figure 28: Multicast in IP-VPN applications.....	104
Figure 29: Dynamic mLDP signaling for IP multicast in VPRN.....	113
Figure 30: MVPN sender-only/receiver-only example.....	115
Figure 31: Multicast VPN traffic flow.....	122
Figure 32: Source PE transit replication and receiver PE per-group SSM extranet mapping (MD-CLI)....	123
Figure 33: Source PE transit replication and receiver PE per-group SSM extranet mapping (classic CLI)..	124
Figure 34: GRT/VRF extranet.....	125
Figure 35: Multicast extranet with per group PIM ASM mapping.....	126
Figure 36: Incongruent multicast and unicast topology for non-overlapping traffic links.....	127
Figure 37: IPv6 MVPN example.....	129
Figure 38: Multicast core diversity.....	130
Figure 39: Logical networks using multi-instance IGP.....	132
Figure 40: NG-MVPN setup via loopback interface.....	133
Figure 41: Intra-AS basic opaque FEC to loopback interface.....	133
Figure 42: Core diversity with parallel NG-MVPN services on parallel IGP instances.....	134
Figure 43: P2MP RSVP-TE core diversity with UFD for UMH redundancy.....	135
Figure 44: Core diversity with parallel NG-MVPN services on parallel IGP instances.....	136

Figure 45: Preferred source selection for multicast source geo-redundancy.....	142
Figure 46: Multicast core diversity.....	145
Figure 47: Inter-AS Option B default MDT setup.....	148
Figure 48: Inter-AS Option C default MDT setup.....	148
Figure 49: Unicast VPN Option B with segmented MPLS.....	150
Figure 50: Unicast VPN Option C with segmented MPLS.....	150
Figure 51: Identical system IP on multiple PEs (Option B).....	152
Figure 52: Non-segmented mLDP PMSI establishment (Option B).....	153
Figure 53: Non-segmented mLDP C-multicast exchange (Option B).....	154
Figure 54: Inter-AS option-C.....	155
Figure 55: Non-segmented mLDP PMSI establishment (Option C).....	156
Figure 56: Non-segmented mLDP C-multicast exchange (Option C).....	157
Figure 57: Bandwidth monitoring.....	164
Figure 58: Example of MVPN upstream PE fast failover.....	167
Figure 59: UMH redundancy for bandwidth monitoring using multiple IOMs.....	168
Figure 60: Bandwidth monitoring to a dedicated IOM.....	170
Figure 61: Redundant monitoring IOMs using a LAG.....	171
Figure 62: OSPF areas.....	198

1 Getting started

1.1 About this guide

This guide describes Layer 3 service functionality provided by 7705 SAR Gen 2 routers and presents examples to configure and implement various protocols and services.

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	

Platform	Collective platform designation
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 Internet Enhanced Service

This chapter provides information about Internet Enhanced Services (IES), the process overview, and implementation notes.

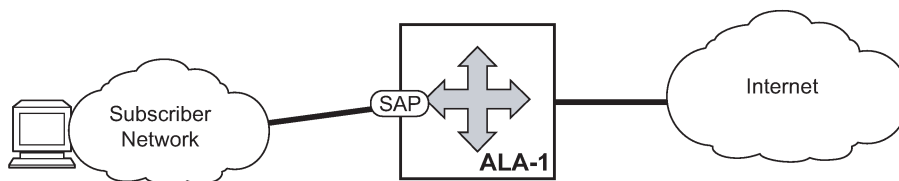
2.1 IES overview

IES is a routed connectivity service where the subscriber communicates with an IP router interface to send and receive Internet traffic. An IES has one or more logical IP routing interfaces, each with a SAP that acts as the access point to the subscriber network.

IES allows IP interfaces to participate in the same routing instance used for service network core routing connectivity. IES services require that the IP addressing scheme used by the subscriber be unique between other provider addressing schemes and potentially the entire Internet. While IES is part of the routing domain, the usable IP address space may be limited. This allows a portion of the service provider address space to be reserved for service IP provisioning, and be administered by a separate, but subordinate address authority.

IP interfaces defined within the context of an IES service must have a SAP associated as the uplink access point to the subscriber network. Multiple IES services are created to segregate subscriber-owned IP interfaces as shown in the following figure.

Figure 1: Internet Enhanced Service



OSSG023

In addition to in-band management connectivity, the IES service supports the following features:

- Multiple IES services are created to separate IP interfaces.
- More than one IES service can be created for a single customer ID.
- More than one IP interface can be created within a single IES service ID. All IP interfaces created within an IES service ID belong to the same customer.

2.2 IES features

This section describes various general service features and any special capabilities or considerations as they relate to IES services.

2.2.1 IP interfaces

IES customer IP interfaces can be configured with most of the same options found on the core IP interfaces. The advanced configuration options supported are:

- VRRP - for IES services with more than one IP interface
- Cflowd
- Secondary IP addresses
- ICMP Options

Configuration options found on core IP interfaces not supported on IES IP interfaces are:

- MPLS forwarding
- NTP broadcast receipt

2.2.1.1 Object grouping and state monitoring

This feature introduces a generic operational group object which associates different service endpoints (pseudowires and SAPs) located in the same or in different service instances. The operational group status is derived from the status of the individual components using specific rules specific to the application using the concept. A number of other service entities, the monitoring objects, can be configured to monitor the operational group status and to perform specific actions as a result of status transitions. For example, if the operational group goes down, the monitoring objects are brought down.

2.2.1.1.1 IES IP interface applicability

This concept is used by an IPv4 IES interface to affect the operational state of the IP interface monitoring the operational group. Individual SAP and spoke SDPs are supported as monitoring objects.

The following rules apply:

- An object can only belong to one group at a time.
- An object that is part of a group cannot monitor the status of a group.
- An object that monitors the status of a group cannot be part of a group.
- An operational group may contain any combination of member types, SAP, or spoke SDPs.
- An operational group may contain members from different VPLS service instances.
- Objects from different services may monitor the operational group.

There are two steps involved in enabling the functionality:

1. Identify a set of objects whose forwarding state should be considered as a whole group, then group them under an operational group using the **oper-group** command.
2. Associate the IP interface to the operational group using the **monitor-oper-group** command.

The status of the operational group is dictated by the status of one or more members according to the following rules:

- The operational group goes down if all the objects in the operational group go down. The operational group comes up if at least one component is up.

- An object in the group is considered down if it is not forwarding traffic in at least one direction. That could be because the operational state is down or the direction is blocked through some validation mechanism.
- If a group is configured but no members are specified yet, then its status is considered up.
- As soon as the first object is configured the status of the operational group is dictated by the status of the provisioned members.

The following configuration shows the operational group g1, the VPLS SAP that is mapped to it and the IP interfaces in IES service 2001 monitoring the operational group g1. This example uses an R-VPLS context.

Operational group g1 has a single SAP (1/1/1:2001) mapped to it and the IP interfaces in the IES service 2001 derive its state from the state of operational group g1.

Example: MD-CLI

In the MD-CLI, the VPLS instance includes the name v1. The IES interface links to the VPLS using the **vpls** command option.

```
[ex:/configure service]
A:admin@node-2# info
  oper-group "g1" {
  }
  ies "2001" {
    customer "1"
    interface "i2001" {
      monitor-oper-group "g1"
      vpls "v1" {
      }
      ipv4 {
        primary {
          address 192.168.1.1
          prefix-length 24
        }
      }
    }
  }
}
vpls "v1" {
  admin-state enable
  service-id 1
  customer "1"
  routed-vpls {
  }
  stp {
    admin-state disable
  }
  sap 1/1/1:2001 {
    oper-group "g1"
    eth-cfm {
      mep md-admin-name "1" ma-admin-name "1" mep-id 1 {
      }
    }
  }
  sap 1/1/2:2001 {
  }
  sap 1/1/3:2001 {
  }
}
```

Example: classic CLI

In the classic CLI, the VPLS instance includes the **allow-ip-int-bind** and the name v1. The IES interface links to the VPLS using the **vpls** command option.

```
A:node-2>config>service# info
-----
oper-group "g1" create
exit
vpls 1 name "v1" customer 1 create
  allow-ip-int-bind
  exit
  stp
    shutdown
  exit
  sap 1/1/1:2001 create
    oper-group "g1"
    eth-cfm
      mep 1 domain 1 association 1 direction down
      ccm-enable
      no shutdown
    exit
  sap 1/1/2:2001 create
    no shutdown
  exit
  sap 1/1/3:2001 create
    no shutdown
  exit
  no shutdown
exit
ies 2001 name "2001" customer 1 create
  shutdown
  interface "i2001" create
    monitor-oper-group "g1"
    address 192.168.1.1/24
    vpls "v1"
    exit
  exit
exit
```

2.2.2 SAPs

This section provides information about SAPs.

2.2.2.1 Encapsulations

The following SAP encapsulations are supported on IES services:

- Ethernet null
- Ethernet dot1q

2.2.2.2 Encapsulation

The packet is encapsulated on an Ethernet pseudowire, which is associated with a pseudowire port on the converged PE, and a spoke SDP on the access PE. The SDP could use an LDP LSP, RSVP LSP,

segment routed tunnel, BGP RFC 8277 tunnel, or LDP over RSVP tunnel. Hash labels are not supported. The SDP may be bound to a port or a LAG, although shaping Vports for pseudowire ports on LAGs in distributed mode is not supported. If an SDP is rerouted, then the corresponding pseudowire ports are brought operationally down. Pseudowire ports are associated with an SDP by configuration.

2.2.2.3 Shaping and bandwidth control

Pseudowire SAPs can be shaped on egress by a Vport on a physical port. The pseudowire SAP egress cannot explicitly declare which Vport to use, but they inherit the Vport used by the PW port egress shaping. The intermediate destination identifier, used for ESM on MPLS pseudowires, is not applicable to VLL, IES, and VPRN pseudowire SAPs.

If a pseudowire port is configured on a LAG, Vport shaping is only supported if the LAG is in link mode.

Per-access node shaping is configured as follows:

1. Configure a Vport per AN under the port (or LAG) to which the SDP corresponding to the pseudowire SAP is bound. Use the **rate** command option in the following context to configure the Vport with an aggregate rate-limit:

- **MD-CLI**

```
configure port ethernet access egress virtual-port aggregate-rate
```

- **classic CLI**

```
configure port ethernet access egress vport agg-rate-limit
```

2. Explicitly assign (by static configuration) a pseudowire port to a Vport. For limiting the total traffic to an AN, all pseudowire ports for an AN-port would refer to the same Vport.

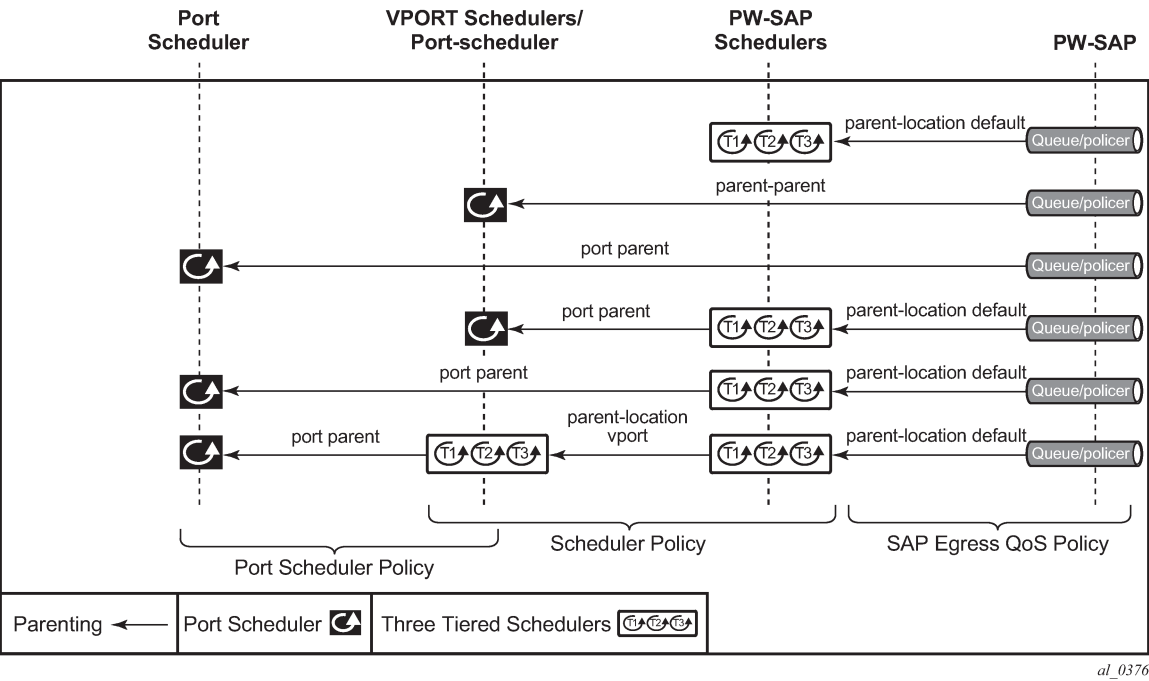
For bandwidth control per pseudowire, the following configuration steps are used:

1. Create multiple Vports under the port to which SDP is bound. Each Vport can be configured with an aggregate rate limit, a scheduler, or a port scheduler.
2. Assign each pseudowire to an AN to a unique Vport shaper (regular IOM/MDA).

To configure the aggregate rate limit or port scheduler under a Vport, PW SAP queues and schedulers must be configured with the **port-parent** command.

To make use of a scheduler under a Vport, PW SAP schedulers must be configured with a **parent** command and the **parent-location vport** under tier 1 of the scheduler policy. The egress hierarchical parenting relationship options are shown in [Figure 2: PW SAP egress scheduling hierarchy options](#). See the *7705 SAR Gen 2 Quality of Service Guide Quality of Service guide* for more information.

Figure 2: PW SAP egress scheduling hierarchy options



2.2.2.4 Lag considerations

PW ports may be bound to Vport schedulers bound to a LAG. However, if the LAG is configured in distributed mode, then bandwidth is shared according to the active LAG members across a single IOM. If the LAG spans multiple IOMs, then it effectively operates in link mode across the IOMs. That is, the full LAG bandwidth is allocated to the LAG members on each IOM. Therefore, the use of a Vport on a distributed mode LAG with a port scheduler on the port or Vport and PW SAPs is explicitly not supported and is not a recommended configuration. It is recommended that port-fair mode is used instead.

2.2.2.5 Last mile packet size adjustment

In the application where pseudowire SAPs are used to apply access QoS for services aggregated from an Ethernet access network, MPLS labels may not be present on the last-mile and link from an access node. In these cases, policers, queues, and H-QoS schedulers should account for packets without MPLS overhead, modeled as “encaps-offset”. Vport and port schedulers behave as per the table below. In the datapath, the actual pseudowire encaps overhead (taking into account the MPLS labels) added to the packet is tracked, and may be applied to the scheduler calculations via the configured packet byte offset.

The rate limit configured for the pseudowire SAP accounts for subscriber or service frame wire rate: without MPLS overhead and including the last mile overhead (unless a packet byte offset is configured).

Table 2: Packet sizes used for pseudowire SAPs summarizes the default packet sizes used at each of the schedulers on the IOM/Ethernet MDA, assuming a 1000 byte customer packet.

Table 2: Packet sizes used for pseudowire SAPs

Type	Size
exp-secondary-shaper	20B preamble + 26 MPLS + 1000B pkt
port-scheduler rate	20B preamble + 1000B pkt
regular queue/policer rate	1000B pkt
vport agg-limit-rate	20B preamble + 1000B pkt
vport port-scheduler rate	20B preamble + 1000B pkt
vport scheduler rate	1000B pkt
vport scheduler to port-scheduler rates	20B preamble + 1000B pkt

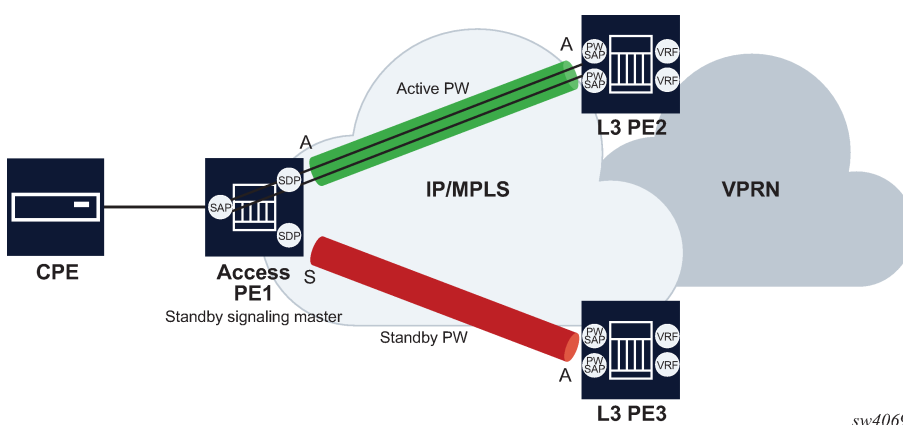
2.2.2.6 Redundancy with pseudowire SAPs

This section describes a mechanism in which one end on a pseudowire (the "master") dictates the active PW selection, which is followed by the other end of the PW (the 'standby'). This mechanism and associated terminology is specified in RFC6870.

Within a chassis, IOM and port based redundancy is based on active/backup LAG. The topology for the base MPLS LSP used by the SDP could be constrained such that it could get re-routed in the aggregation network, but would always appear on the LAG ports on the Layer 3 PE. In the case that the tunnel is re-routed to a different port, the MPLS pseudowire SAPs would be brought down.

To provide Layer 3 PE redundancy, dual homing of the access PE into separate Layer 3 PEs using active/standby pseudowire status is supported. This is shown in [Figure 3: Dual homing into multiple Layer 3 PEs](#).

Figure 3: Dual homing into multiple Layer 3 PEs



Dual homing operates in a similar manner to Spoke-SDP termination on IES/VPRN. [Figure 3: Dual homing into multiple Layer 3 PEs](#) displays the access PE is dual-homed to the Layer 3 PEs using two spoke-SDPs. Use the following command to configure the endpoint in the access PE to be the master from a pseudowire redundancy perspective:

- **MD-CLI**

```
configure service epipe endpoint standby-signaling master
```

- **classic CLI**

```
configure service epipe endpoint standby-signaling-master
```

The access PE picks one of the spoke SDPs to make active, and one to make standby, based on the local configuration of primary or spoke SDP precedence.

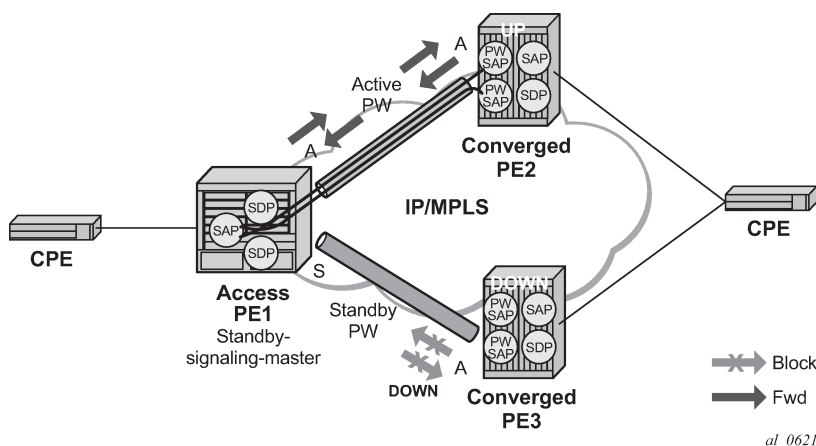
The pseudowire port at the Layer 3 PE behaves as a standby from the perspective of pseudowire status signaling. That is, if its peer signals "PW FWD standby (0x20)" status bit for the specific Spoke-SDP and the local configuration does not allow this bit to be ignored, the PE takes the pseudowire port to a local operationally down state. This is consistent with the Spoke-SDP behavior for the case of Spoke-SDP termination on IES/VPRN.

As a consequence, all of the pseudowire SAPs bound to the pseudowire port are taken down, which causes the corresponding IES or VPRN interface to go to a local operationally down state and therefore stops forwarding packets toward this pseudowire port.

Conversely, the formerly standby pseudowire is made active and then the corresponding pseudowire port on the second Layer 3 PE is taken locally operationally up. Therefore, all of the pseudowire SAPs bound to the pseudowire port are brought up, which causes the corresponding IES or VPRN interface to go to a local operationally up state allowing forwarding of packets toward this pseudowire port.

For VLLs, a PW port always behaves as a standby from the perspective of PW redundancy. This is because the PW port is taken locally operationally down if any non-zero PW status (including a PW Preferential Forwarding status of standby) is received. Master-standby PW redundancy mechanisms for dual homing of the access PE into separate converged PE using active or standby PW status are supported as shown in [Figure 4: Master-standby PW redundancy](#). However, this is only applicable to VLL services consisting of only SAPs, PW-SAPs, or spoke-SDPs. Dual-homing redundancy, taking into account the status of the PW SAP, is not supported where a PBB tunnel between the two converged PE exists in the Epipe VLL service.

Figure 4: Master-standby PW redundancy



As in the existing implementation, standby signaling is configured to master on the spoke SDP at the access PE. However, explicit configuration of standby signaling to the standby on the PW port is not required, as this is the default behavior.

The forwarding behavior is the same as when standby-signaling is configured to standby for Epipe spoke SDPs. That is, when enabled, if a PW Forwarding Standby (0x20) LDP status message is received for the P1111111W, then the transmit direction is blocked for the PW port. All PW SAPs bound to the corresponding PW port are treated from a SAP OAM perspective in the same manner as a fault on the service, such as an SDP-binding down or remote SAP down.

PW redundancy with multiple active/standby PW ports or PW SAPs bound to the same Ethernet SAP in the converged PE is not supported. The independent mode of operation for PW redundancy is not also supported for a PW port.

2.2.2.7 Operational group support for PW ports

A PW port state may be linked to the state of an operational group, such that if the operational group goes down, the SDP binding for the PW port also goes operationally down, and therefore the corresponding PW status bit signaled (0x00000001 - Pseudowire Not Forwarding). If a status of 0x00000001 is signaled for a currently active PW, and active/standby dual homing is in use then the access PE fails over to the standby PW to the standby converged PE.

This is achieved by linking an SDP binding to an operational group for PW SAPs belonging to any supported service types (including those with group interfaces) bound to that PW port, such as IES, VPRN, or Epipe VLL.

Example: Associate an operational group at the SDP binding level (MD-CLI)

```
[ex:/configure service]
A:admin@node-2# info
    sdp 1 {
    }
    pw-port {
        binding-port 1/1/1
    }
}

[ex:/configure pw-port 1 sdp 1]
A:admin@node-2# info
    admin-state enable
    vc-id 11
    monitor-oper-group "test-oper-grp"
```

Example: Associate an operational group at the SDP binding level (classic CLI)

```
A:node-2>config>service# info
-----
    sdp 1 create
        no shutdown
        binding
            port 1/1/1
            pw-port 1 vc-id 11 create
                no shutdown
                monitor-oper-group "test-oper-grp"
            exit
        exit
    exit
```

The **monitor-oper-group** command specifies the operational group to be monitored by the PW-Port under which it is configured. The operational group name must be already configured under the **configure service** context before its name is referenced in this command.

The following illustrates how a PW port can track the status of VPRN uplinks using **monitor-oper-group**. Uplinks in a VPRN may be monitored using a BFD session on the network facing IP interfaces in a VPRN or on the network IP interfaces supporting the uplinks.

Example: Configure an operational group to monitor the BFD session (MD-CLI)

```
[ex:/configure service]
A:admin@node-2# info
  oper-group "test-oper-grp" {
    bfd-liveness {
      router-instance "105"
      interface-name "vprn-if"
      dest-ip 10.0.0.20
    }
  }
```

Example: Configure an operational group to monitor the BFD session (classic CLI)

```
A:node-2>config>service# info
-----
  oper-group "test-oper-grp" create
    bfd-enable interface "vprn-if" dest-ip 10.0.0.20 service 105
```

Alternatively, the state of network interfaces can be monitored as follows.

Example: Configure the monitoring of network interfaces (MD-CLI)

```
[ex:/configure service]
A:admin@node-2# info
  oper-group "test-oper-grp" {
    bfd-liveness {
      interface-name "network-if"
      dest-ip 10.0.1.20
    }
  }
```

Example: Configure the monitoring of network interfaces (classic CLI)

```
A:node-2>config>service# info
-----
  oper-group "test-oper-grp" create
    bfd-enable interface "network-if" dest-ip 10.0.1.20
```

The PW port is then configured with **monitor-oper-group** as follows.

Example: Configure the PW port with the monitor-oper-group option (MD-CLI)

```
[ex:/configure service]
A:admin@node-2# info
  sdp 1 {
  }
  pw-port {
    binding-port 1/1/2
  }
}

[ex:/configure pw-port 100 sdp 1]
A:admin@node-2# info
  vc-id 25
```

```
monitor-oper-group "test-oper-grp"
```

Example: Configure the PW port with the monitor-oper-group option (classic CLI)

```
A:node-2>config>service>sdp>binding# info
-----
    port 1/1/2
    pw-port 100 vc-id 25 create
        shutdown
        monitor-oper-group "test-oper-grp"
    exit
-----
```

2.2.3 Routing protocols

The IES IP interfaces are restricted as to the routing protocols that can be defined on the interface based on the fact that the customer has a different routing domain for this service. The IES IP interfaces support the following routing protocols:

- RIP
- OSPF
- IS-IS
- BGP
- IGMP
- PIM



Note: The SAP for the IES IP interface is created at the IES service level, but the routing protocols for the IES IP interface are configured at the routing protocol level for the main router instance.

2.2.3.1 CPE connectivity check

Static routes are used within many IES services. Unlike dynamic routing protocols, there is no way to change the state of routes based on availability information for the associated CPE. CPE connectivity check adds flexibility so that unavailable destinations are removed from the service provider routing tables dynamically and minimize wasted bandwidth.

The availability of the far-end static route is monitored through periodic polling. The polling period is configured. If the poll fails a specified number of sequential polls, the static route is marked as inactive.

An ICMP ping mechanism is used to test the connectivity. If the connectivity check fails and the static route is deactivated, the router continues to send polls and reactivate any routes that are restored.

2.2.4 QoS policies

When applied to IES services, service ingress QoS policies only create the unicast queues defined in the policy. The multipoint queues are not created on the service. With IES services, service egress QoS policies function as with other services where the class-based queues are created as defined in the policy. Both Layer 2 or Layer 3 criteria can be used in the QoS policies for traffic classification in an IES.

2.2.5 Filter policies

Only IP filter policies can be applied to IES services.

2.2.6 MPLS EL and hash label

The router supports both the MPLS entropy label (EL) (RFC 6790) and the Flow Aware Transport (FAT) label, known as the hash label (RFC 6391). 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 EL is supported for Epipe spoke-SDP termination on IES interfaces. To configure insertion of the EL on a spoke SDP that terminates on the service, use the **entropy-label** command in the **spoke-sdp** context of a specific interface.

The hash label is also supported for Epipe spoke-SDP termination on IES services. Configure it using the **hash-label** command in the **spoke-sdp** context for an IES interface context.

Either the hash label or the EL can be configured on one object, but not both.

2.2.7 Spoke SDPs

Distributed services use service destination points (SDPs) to direct traffic to another router through service tunnels. SDPs are created on each participating router and then bound to a specific service. SDP can be created as either GRE or MPLS. See the *7705 SAR Gen 2 Services Overview Guide* for information about configuring SDPs.

This feature provides the ability to cross-connect traffic entering on a spoke SDP, used for Layer 2 services (VLLs or VPLS), on to an IES or VPRN service. From a logical point of view, the spoke SDP entering on a network port is cross-connected to the Layer 3 service as if it entered by a service SAP. The main exception to this is traffic entering the Layer 3 service by a spoke SDP is handled with network QoS policies and not access QoS policies.

[Figure 5: SDP-ID and VC label service identifiers](#) depicts traffic terminating on a specific IES or VPRN service that is identified by the SDP-ID and VC label present in the service packet.

Figure 5: SDP-ID and VC label service identifiers

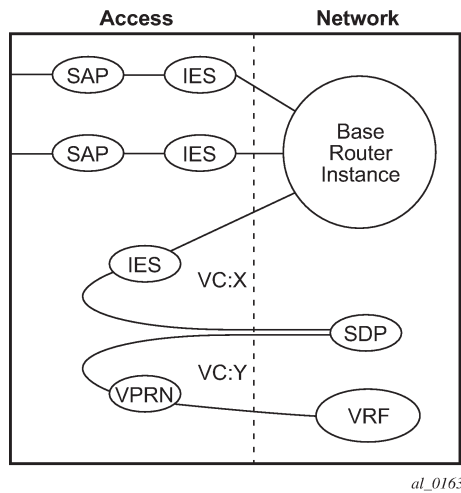


Figure 6: IES spoke-SDP termination

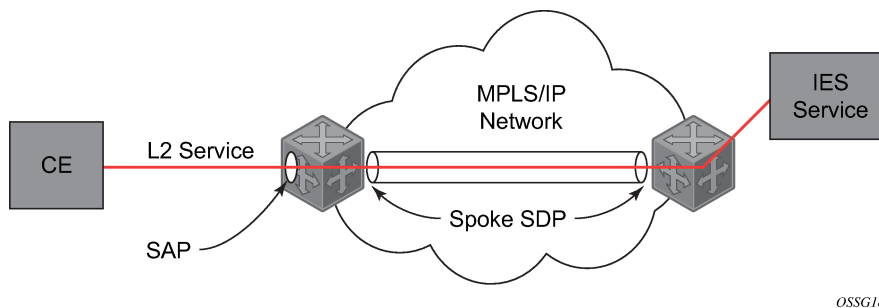


Figure 6: IES spoke-SDP termination depicts a spoke SDP terminating directly into a Layer 3 service interface (IES or VPRN) at one end, and a Layer 2 service (Epipe or VPLS) at the other. There is no special configuration required on the Layer 2 service.

Spoke SDPs created with **vc-type ether** (the default) are compatible with Epipe and VPLS services, as well as with other IES/VPRN interfaces.

If the MPLS network uses LDP signaling, then in order for a spoke SDP to function, the LDP binding MTUs at each end must match. For a Layer 2 service, the MTU of the local binding is 14 octets less than the configured service MTU (such as, binding MTU = service MTU - 14). For an IES or VPRN interface, the binding MTU is equal to either the configured **ip-mtu** of the interface, or the SDP's **path-mtu** minus 14, whichever is lower. Use the following command to find the local and remote MTUs of all bindings.

```
show router ldp bindings
```

All routing protocols that are supported by IES/VPRN are supported for spoke-SDP termination.

See "VCCV BFD support for VLL, spoke SDP termination on IES and VPRN, and VPLS services" in the *7705 SAR Gen 2 Layer 2 Services and EVPN Guide* for information about using VCCV BFD in spoke-SDP termination.



Note: Spoke-SDP termination of Ipipe VLLs on IES is not supported in System Profile B. Use the following command to determine if Ipipes are currently bound to an IES interface before configuring profile B.

```
show router ldp bindings services
```

2.2.8 DHCP

This section describes DHCP feature support for IES services.

2.2.8.1 DHCP client for IES

The 7705 SAR Gen 2 supports IES interfaces configured with a DHCP client. When the node operates as a DHCP client, it learns the IP address of the interface via dynamic IP address assignment.

The DHCP client implementation for IES is identical to that of the Base router context. See "DHCP client" in the *7705 SAR Gen 2 Router Configuration Guide* for more information.



Note: When DHCP client is enabled on a SAP interface, a second SAP interface with or without a DHCP client must be configured (in either the administratively enabled or disabled state) for the DHCP client to become operational on the original SAP interface.

2.2.8.2 DHCP relay for IES

The 7705 SAR Gen 2 supports the DHCP relay agent for IES interfaces. When the DHCP relay agent is enabled, broadcast DHCP request messages are intercepted and forwarded unicast to a configured DHCP server.

The DHCP relay agent implementation for IES is identical to that of the Base router context. See "DHCP relay" in the *7705 SAR Gen 2 Router Configuration Guide* for more information.

2.2.8.3 Local DHCP server for IES

The 7705 SAR Gen 2 supports a local DHCP server for IES interfaces. When the DHCP relay agent is enabled, broadcast DHCP request messages are intercepted and can be forwarded unicast to a locally configured DHCP server.

See "DHCP relay" in the *7705 SAR Gen 2 Router Configuration Guide* for more information about the local DHCP server function.

See "DHCP persistence" in the *7705 SAR Gen 2 Basic System Configuration Guide* for more information about DHCP server persistency.

2.3 Configuring an IES service with CLI

This section provides information to configure IES services using the CLI.

2.3.1 Basic configuration

The most basic IES service configuration has the following entities:

- customer ID (see the *7705 SAR Gen 2 Services Overview Guide* for more information)
- an interface to create and maintain IP routing interfaces within IES service ID
- a SAP on the interface specifying the access port and encapsulation values

The following example shows the configuration of an IES service.

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2# info
  ies "1000" {
    admin-state enable
    description "to internet"
    customer "1"
    vpn-id 1000
    interface "to-web" {
      sap 1/1/5:0.* {
      }
      ipv4 {
        primary {
          address 10.1.1.1
          prefix-length 24
        }
      }
    }
  }
}
```

Example: classic CLI

```
A:node-2>config>service# info
-----
  ies 1000 name "1000" customer 1 vpn 1000 create
    description "to internet"
    interface "to-web" create
      address 10.1.1.1/24
      sap 1/1/5:0.* create
    exit
  exit
no shutdown
exit
-----
```

2.3.2 Common configuration tasks

About this task

This section provides a brief overview of the tasks that must be performed to configure IES services and provides the CLI commands.

Perform the following tasks to configure IES services.

Procedure

- Step 1.** Associate an IES service with a customer ID.
- Step 2.** Associate customer ID with the service.
- Step 3.** Assign an IP address.
- Step 4.** Optional: Create a subscriber interface.
- Step 5.** Create an interface.
- Step 6.** Define SAP command options on the interface.
 - a. Select nodes and ports.
 - b. Optional: Select QoS policies other than the default (configured in the **configure qos** context).
 - c. Optional: Select filter policies (configured in the **configure filter** context).
 - d. Optional: Select accounting policy (configured in the **configure log** context).
- Step 7.** Enable service.

2.3.3 Configuring IES components

Use the configuration information that follows to configure IES components.

2.3.3.1 Configuring an IES service

The following example displays a basic IES service configuration.

Example: IES service configuration for the 7705 SAR Gen 2 (MD-CLI)

```
[ex:/configure service]
A:admin@node-2# info
  ies "1001" {
    admin-state enable
    description "to-internet"
    customer "1"
    vpn-id 1001
  }
```

Example: IES service configuration for the 7705 SAR Gen 2 (classic CLI)

```
A:node-2>config>service# info
-----
  ies 1001 name "1001" customer 1 vpn 1001 create
    description "to-internet"
    no shutdown
  exit
-----
```

2.3.3.2 Configuring IES subscriber interfaces



Note: This section applies to the 7705 SAR Gen 2 only.

Subscriber interfaces operate only with basic (or enhanced) subscriber management. At the very least, a host, either statically configured or dynamically learned by DHCP must be present in order for the interface to be useful.

The following example displays a subscriber interface configuration for the 7705 SAR Gen 2.

Example: MD-CLI

```
[ex:/configure service ies "11" subscriber-interface "test1"]
A:admin@node-2# info
  ipv4 {
    address 192.168.140.1 {
      prefix-length 24
    }
  }
  group-interface "abc-if" {
    sap 1/1/19:0 {
      ingress {
        qos {
          sap-ingress {
            policy-name "2"
          }
        }
        filter {
          ip "10"
        }
      }
      static-host {
        ipv4 192.168.145.100 mac 00:01:00:00:00:01 {
        }
      }
    }
  }
}
```

Example: classic CLI

```
A:node-2>config>service>ies>sub-if# info
-----
      address 192.168.140.1/24
      group-interface "abc-if" create
      sap 1/1/19:0 create
      ingress
        qos 2
        filter ip 10
      exit
      static-host ip 192.168.145.100 mac 00:01:00:00:00:01 create
      exit
    exit
  exit
-----
```

2.3.3.3 Configuring IES interfaces

The following examples show the configuration of IES interfaces.

Example: IES configuration for the 7705 SAR Gen 2 (MD-CLI)

```
[ex:/configure service ies "2" interface "test2"]
A:admin@node-2# info
  sap 1/1/10:0.* {
    ingress {
      qos {
        sap-ingress {
          policy-name "100"
        }
      }
    }
    egress {
      qos {
        scheduler-policy {
          policy-name "SLA1"
        }
      }
    }
  }
  ipv4 {
    primary {
      address 10.1.1.1
      prefix-length 24
    }
    vrrp 1 {
      authentication-key "McTNkSePNJMVfysxyZa4y4D+iaD7lJ4= hash2"
    }
  }
}
```

Example: IES configuration for the 7705 SAR Gen 2 (classic CLI)

```
A:node-2>config>service>ies>if$ info
-----
      address 10.1.1.1/24
      vrrp 1 owner
      authentication-key "McTNkSePNJMVfysxyZa4y4D+iaD7lJ4=" hash2
      exit
      sap 1/1/10:0.* create
      ingress
      qos 100
      exit
      egress
      scheduler-policy "SLA1"
      exit
      exit
-----
```

2.3.3.4 Configuring a spoke SDP

The following example shows a spoke-SDP configuration for the 7705 SAR Gen 2.

Example: MD-CLI

```
[ex:/configure service ies "4"]
A:admin@node-2# info
  admin-state enable
  description "to internet"
  customer "1"
  interface "spokeSDP-test" {
    spoke-sdp 2:100 {
      egress {
        filter {
          ip "10"
        }
      }
    }
  }
}
```

Example: classic CLI

```
A:node-2>config>service>ies# info
-----
      description "to internet"
      interface "spokeSDP-test" create
        spoke-sdp 2:100 create
          egress
            filter ip 10
          exit
        exit
      exit
    exit
  no shutdown
-----
```

2.3.3.5 Configuring a SAP

A service access point (SAP) is a combination of a port and encapsulation command options that identifies the SAP on the interface and within the router. Each SAP must be unique within a router.

When configuring IES SAP command options on a 7705 SAR Gen 2, a default QoS policy is applied to each ingress and egress SAP. Additional QoS policies and scheduler policies must be configured in the **configure qos** context. Filter policies are configured in the **configure filter** context and must be explicitly applied to a SAP. There are no default filter policies.

Example: IES SAP configuration for the 7705 SAR Gen 2 (MD-CLI)

```
[ex:/configure service ies "2"]
A:admin@node-2# info
  customer "1"
  interface "test2" {
    sap 5/1/3.1:0 {
      ingress {
        qos {
          sap-ingress {
            policy-name "101"
          }
        }
      }
    }
    egress {
      qos {
        sap-egress {
```

```

        policy-name "1010"
      }
      scheduler-policy {
        policy-name "alpha"
      }
    }
  }
  ipv4 {
    primary {
      address 10.10.36.2
      prefix-length 24
    }
  }
}

```

Example: IES SAP configuration for the 7705 SAR Gen 2 (classic CLI)

```

A:node-2>config>service>ies>if# info
-----
      address 10.10.36.2/24
      sap 5/1/3.1:0 create
      ingress
        qos 101
      exit
      egress
        scheduler-policy "alpha"
        qos 1010
      exit
    exit
  -----

```

2.3.3.6 Configuring VRRP

VRRP can be configured in either an owner or non-owner mode. The owner is the VRRP router whose virtual router IP address is the same as the real interface IP address. This is the router that responds to packets addressed to one of the IP addresses for ICMP pings, TCP connections, and so on. All other virtual router instances participating in this message domain must have the same VRID configured and cannot be configured as owner.

See the *7705 SAR Gen 2 Router Configuration Guide* for more information about VRRP.

The following example shows the VRRP on an IES interface configuration.

Example: MD-CLI

```

[ex:/configure service ies "16"]
A:admin@node-2# info
  customer "1"
  interface "test16" {
    ipv4 {
      primary {
        address 10.10.36.2
        prefix-length 24
      }
      vrrp 2 {
        authentication-key "McTNkSePNJMVfysxyZa4y1fsBIad0g= hash2"
        backup [10.10.36.2]
        owner true
      }
    }
  }

```

```
    }
}
```

Example: classic CLI

```
A:node-2>config>service>ies>if$ info
-----
        address 10.10.36.2/24
        vrrp 2 owner
            backup 10.10.36.2
            authentication-key "McTNkSePNJMVFysxyZa4y1fsBIiad0g=" hash2
        exit
-----
```

2.3.3.7 Configuring IPsec

The following example shows an IES service with IPsec configured for the 7705 SAR Gen 2.

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2# info
  ies "100" {
    admin-state enable
    customer "1"
    interface "ipsec-public" {
      admin-state enable
      sap ipsec-1.public:1 {
      }
      ipv4 {
        primary {
          address 10.10.10.1
          prefix-length 24
        }
      }
    }
  }
}
```

Example: classic CLI

```
A:node-2>config# info
-----
...
  service
    ies 100 customer 1 create
      interface "ipsec-public" create
        address 10.10.10.1/24
        sap ipsec-1.public:1 create
      exit
    exit
  no shutdown
  exit
exit
...
-----
```

2.3.3.8 IGMP host tracking

The following example shows an IES service with IGMP host tracking configured.

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2# info
...
  ies "25" {
    admin-state enable
    customer "1"
    interface "ip_if_4" {
      ip-mtu 9000
      tos-marking-state untrusted
      sap lag-64:64 {
      }
      hold-time {
        ipv4 {
          down {
            seconds 1200
          }
        }
      }
      ipv4 {
        allow-directed-broadcasts true
        local-dhcp-server "server 1"
        urpf-check {
        }
        primary {
          address 10.64.64.64
          prefix-length 24
        }
        secondary 10.3.4.5 {
          prefix-length 24
        }
        neighbor-discovery {
          local-proxy-arp true
          remote-proxy-arp true
          proxy-arp-policy ["treetrace-1"]
        }
        dhcp {
          admin-state enable
          description "server 1"
          server [192.168.17.1]
        }
      }
    }
    igmp-host-tracking {
      expiry-time 65535
    }
  }
```

Example: classic CLI

```
A:node-2>config>service# info
-----
...
  ies 25 name "25" customer 1 create
    shutdown
    igmp-host-tracking
      expiry-time 65535
```

```

        no shutdown
    exit
    interface "ip_if_4" create
        hold-time
        down ip 1200
    exit
    address 10.64.64.64/24
    secondary 10.3.4.5/24
    allow-directed-broadcasts
    tos-marking-state trusted
    dhcp
        shutdown
        description "server 1"
        server 192.168.17.1
    exit
    ip-mtu 9000
    local-dhcp-server "server 1"
    local-proxy-arp
    proxy-arp-policy tree-trace-1
    remote-proxy-arp
    sap lag-64:64 create
    exit
    urpf-check
    exit
exit
exit

```

2.4 Service management tasks

This section describes the IES service management tasks.

2.4.1 Modifying IES service command options

Existing IES service command options in the CLI or NMS can be modified, added, removed, enabled or disabled. Changes you make are applied immediately to all services.

To display a list of customer IDs, use the **show service customer** command. Enter the command options (such as description, SAP information and SDP information), and then enter the new information.

Example: Modified service configuration for the 7705 SAR Gen 2 (MD-CLI)

```

[ex:/configure service]
A:admin@node-2# info
  ies "1000" {
    admin-state enable
    description "This is a new description"
    customer "1"
    vpn-id 1000
    interface "to-web" {
      mac 00:dc:98:1d:00:00
      sap 22/1/50:0 {
      }
    }
    ipv4 {
      allow-directed-broadcasts true
      primary {
        address 10.1.1.1
        prefix-length 24
      }
    }
  }

```

```

    }
  }
}

```

Example: Modified service configuration for the 7705 SAR Gen 2 (classic CLI)

```

A:node-2>config>service# info
-----
    ies 1000 name "1000" customer 1 vpn 1000 create
      description "This is a new description"
      interface "to-web" create
        address 10.1.1.1/24
        mac 00:dc:98:1d:00:00
        allow-directed-broadcasts
        sap 22/1/50:0 create
        exit
      exit
    no shutdown
  exit
-----

```

2.4.2 Deleting a spoke SDP

The following example shows how to delete a spoke SDP (10.100) from the interface configuration (spoke-SDP-test interface) for the 7705 SAR Gen 2.

Example: Delete a spoke SDP (MD-CLI)

```

[ex:/configure service ies "14" interface "spokeSDP-test"]
A:admin@node-2# delete spoke-sdp 10:100

```

Example: Delete a spoke SDP (classic CLI)

In classic CLI, the service interface must be shutdown to delete the spoke SDP. This cleans up the associated VC labels.

```

*A:node-2>config>service>ies# shutdown
*A:node-2>config>service>ies# interface "spokeSDP-test"
*A:node-2>config>service>ies>if# shutdown
*A:node-2>config>service>ies>if# spoke-sdp 10:100
*A:node-2>config>service>ies>if>spoke-sdp# shutdown
*A:node-2>config>service>ies>if>spoke-sdp# exit
*A:node-2>config>service>ies>if# no spoke-sdp 10:100

```

2.4.3 Deleting an IES service

The following example shows the deletion of an IES service.

Example: Delete an IES service (MD-CLI)

```

[ex:/configure service]
A:admin@node-2# delete ies 13

```

Example: Delete an IES service (classic CLI)

In classic CLI, an IES service cannot be deleted until SAPs and interfaces are shut down and deleted and the service is shutdown on the service level.

```
*A:node-2>config>service>ies# interface "test3"
*A:node-2>config>service>ies>if# sap 1/1/7:0.*
*A:node-2>config>service>ies>if>sap# shutdown
*A:node-2>config>service>ies>if>sap# exit
*A:node-2>config>service>ies>if# no sap 1/1/7:0.*
*A:node-2>config>service>ies>if# shutdown
*A:node-2>config>service>ies>if# exit
*A:node-2>config>service>ies# no interface "test3"
*A:node-2>config>service>ies# shutdown
*A:node-2>config>service>ies# exit
*A:node-2>config>service# no ies 13
```

2.4.4 Disabling an IES service

The following example shows the disabling of an IES service. You can disable an IES service without deleting the service command options.

Example: Disabling of an IES service (MD-CLI)

```
[ex:/configure service ies "12"]
A:admin@node-2# admin-state disable

[ex:/configure service ies "12"]
A:admin@node-2# info
  admin-state disable
  customer "1"
  interface "test2" {
  ...
```

Example: Disabling of an IES service (classic CLI)

```
A:node-2>config>service>ies# shutdown
A:node-2>config>service>ies# info
-----
      shutdown
      interface "test2" create
  ...
```

2.4.5 Re-enabling an IES service

The following example shows how to re-enable an IES service that was administratively disabled.

Example: Re-enabling an IES service (MD-CLI)

```
[ex:/configure service ies "12"]
A:admin@node-2# admin-state enable

[ex:/configure service ies "12"]
A:admin@node-2# info
  admin-state enable
  customer "1"
```

```
interface "test2" {  
...  
}
```

Example: Re-enabling an IES service (classic CLI)

```
A:node-2>config>service>ies# no shutdown  
A:node-2>config>service>ies# info  
-----  
interface "test2" create  
    address 10.1.1.1/24  
...  
}
```

3 VPRN service

This chapter provides information about the VPRN service and implementation.

3.1 VPRN service overview

RFC 2547b is an extension to the original RFC 2547, *BGP/MPLS VPNs*, which details a method of distributing routing information using BGP and MPLS forwarding data to provide a Layer 3 Virtual Private Network (VPN) service to end customers.

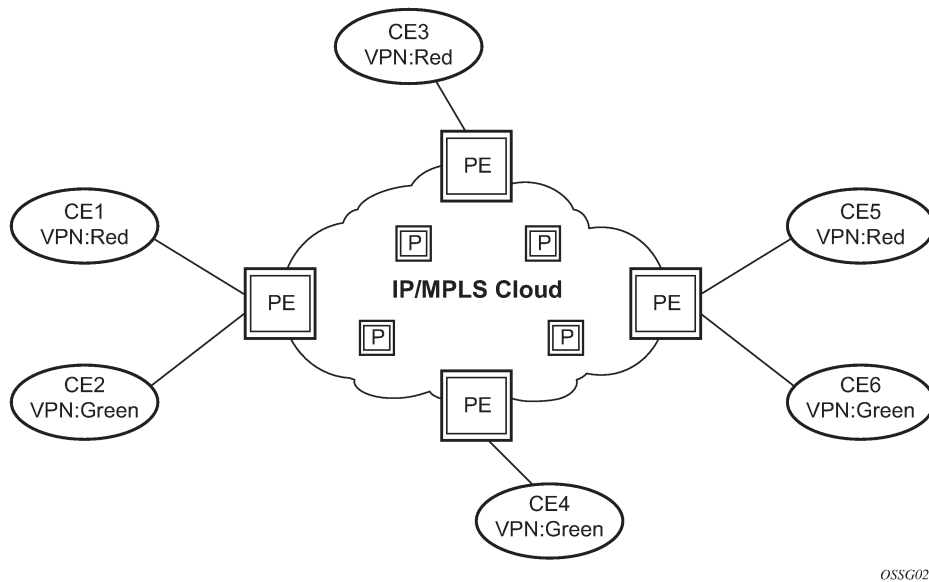
Each Virtual Private Routed Network (VPRN) consists of a set of customer sites connected to one or more PE routers. Each associated PE router maintains a separate IP forwarding table for each VPRN. Additionally, the PE routers exchange the routing information configured or learned from all customer sites via MP-BGP peering. Each route exchanged via the MP-BGP protocol includes a Route Distinguisher (RD), which identifies the VPRN association and handles the possibility of IP address overlap.

The service provider uses BGP to exchange the routes of a particular VPN among the PE routers that are attached to that VPN. This is done in a way which ensures that routes from different VPNs remain distinct and separate, even if two VPNs have an overlapping address space. The PE routers peer with locally connected CE routers and exchange routes with other PE routers to provide end-to-end connectivity between CEs belonging to a specific VPN. Because the CE routers do not peer with each other there is no overlay visible to the CEs.

When BGP distributes a VPN route it also distributes an MPLS label for that route. On an SR series router, the method of allocating a label to a VPN route depends on the VPRN label mode and the configuration of the VRF export policy. SR series routers support three label allocation methods: label per VRF, label per next hop, and label per prefix.

Before a customer data packet travels across the service provider's backbone, it is encapsulated with the MPLS label that corresponds, in the customer's VPN, to the route which best matches the packet's destination address. The MPLS packet is further encapsulated with one or additional MPLS labels or GRE tunnel header so that it gets tunneled across the backbone to the correct PE router. Each route exchanged by the MP-BGP protocol includes a route distinguisher (RD), which identifies the VPRN association. Thus the backbone core routers do not need to know the VPN routes. [Figure 7: Virtual Private Routed Network](#) displays a VPRN network diagram example.

Figure 7: Virtual Private Routed Network



3.1.1 Routing prerequisites

RFC 4364 requires the following features:

- multiprotocol extensions to BGP
- extended BGP community support
- BGP capability negotiation

Tunneling protocol options are as follows:

- Label Distribution Protocol (LDP)
- MPLS RSVP-TE tunnels
- Generic Router Encapsulation (GRE) tunnels
- BGP route tunnel (RFC 8277)

3.1.2 Core MP-BGP support

BGP is used with BGP extensions mentioned in [Routing prerequisites](#) to distribute VPRN routing information across the service provider's network.

BGP was initially designed to distribute IPv4 routing information. Therefore, multiprotocol extensions and the use of a VPN-IP address were created to extend BGP's ability to carry overlapping routing information. A VPN-IPv4 address is a 12-byte value consisting of the 8-byte route distinguisher (RD) and the 4-byte IPv4 IP address prefix. A VPN-IPv6 address is a 24-byte value consisting of the 8-byte RD and 16-byte IPv6 address prefix. Service providers typically assign one or a small number of RDs per VPN service network-wide.

3.1.3 Route distinguishers

The route distinguisher (RD) is an 8-byte value consisting of two major fields: the type field and the value field, as shown in the following figure. The type field determines how the value field should be interpreted. Three type values are supported as defined in the Internet draft.

Figure 8: Route distinguisher



The three type values are:

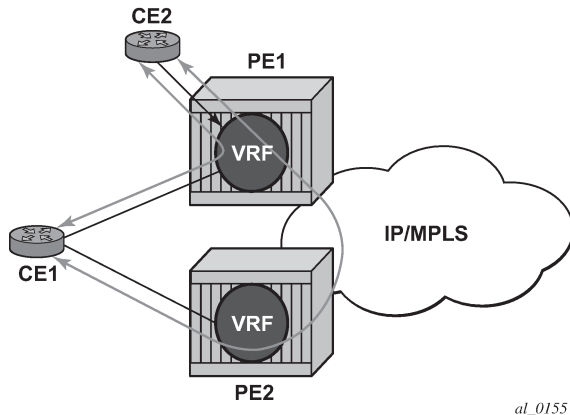
- Type 0: Value Field — Administrator subfield (2 bytes)
Assigned number subfield (4 bytes)
The administrator field must contain an ASN (using private ASNs is discouraged). The Assigned field contains a number assigned by the service provider.
- Type 1: Value Field — Administrator subfield (4 bytes)
Assigned number subfield (2 bytes)
The administrator field must contain an IP address (using private IP address space is discouraged). The Assigned field contains a number assigned by the service provider.
- Type 2: Value Field — Administrator subfield (4 bytes)
Assigned number subfield (2 bytes)
The administrator field must contain a 4-byte ASN (using private ASNs is discouraged). The Assigned field contains a number assigned by the service provider.

3.1.3.1 eiBGP load balancing

eiBGP load balancing allows a route to have multiple next hops of different types, using both IPv4 next hops and MPLS LSPs simultaneously.

Figure 9: Basic eiBGP topology displays a basic topology that could use eiBGP load balancing. In this topology CE1 is dual homed and therefore reachable by two separate PE routers. CE 2 (a site in the same VPRN) is also attached to PE1. With eiBGP load balancing, PE1 uses its own local IPv4 next hop as well as the route advertised by MP-BGP, by PE2.

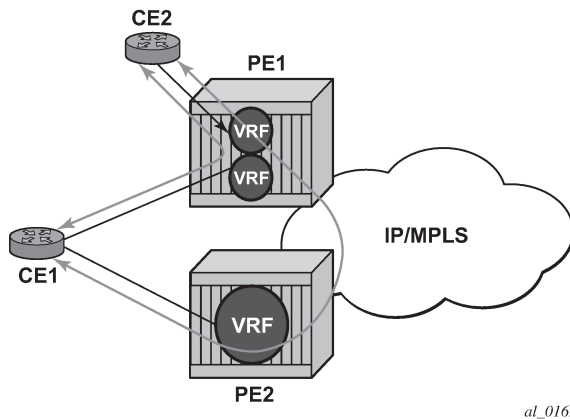
Figure 9: Basic eiBGP topology



Another example displayed in [Figure 10: Extranet load balancing](#) shows an extra net VPRN (VRF). The traffic ingressing the PE that should be load balanced is part of a second VPRN and the route over which the load balancing is to occur is part of a separate VPRN instance and are leaked into the second VPRN by route policies.

Here, both routes can have a source protocol of VPN-IPv4 but one still has an IPv4 next hop and the other can have a VPN-IPv4 next hop pointing out a network interface. Traffic is still load balanced (if eiBGP is enabled) as if only a single VRF was involved.

Figure 10: Extranet load balancing



Traffic is load balanced across both the IPv4 and VPN-IPv4 next hops. This helps to use all available bandwidth to reach a dual-homed VPRN.

3.1.4 Route reflector

The use of Route Reflectors is supported in the service provider core. Multiple sets of route reflectors can be used for different types of BGP routes, including IPv4 and VPN-IPv4 as well as multicast and IPv6 (multicast and IPv6 apply to the 7750 SR only).

3.1.5 CE to PE route exchange

Routing information between the Customer Edge (CE) and Provider Edge (PE) can be exchanged by the following methods:

- static routes
- eBGP
- RIP
- OSPF
- OSPF3

Each protocol provides controls to limit the number of routes learned from each CE router.

3.1.5.1 Route redistribution

Routing information learned from the CE-to-PE routing protocols and configured static routes should be injected in the associated local VPN routing or forwarding (VRF). In the case of dynamic routing protocols, there may be protocol-specific route policies that modify or reject specific routes before they are injected into the local VRF.

Route redistribution from the local VRF to CE-to-PE routing protocols is controlled via the route policies in each routing protocol instance, in the same manner that is used by the Base router instance.

The advertisement or redistribution of routing information from the local VRF to or from the MP-BGP instance is specified per VRF and is controlled by VRF route target associations or by VRF route policies.

VPN-IP routes imported into a VPRN have the protocol type **bgp-vpn** to denote that it is a VPRN route. Use the following command to configure this within the route policy match criteria:

- **MD-CLI**

```
configure policy-options policy-statement entry from protocol name
```

- **classic CLI**

```
configure router policy-options policy-statement entry from protocol
```

3.1.5.2 CPE connectivity check

Static routes are used within many IES services and VPRN services. Unlike dynamic routing protocols, there is no way to change the state of routes based on availability information for the associated CPE. CPE connectivity check adds flexibility so that unavailable destinations are removed from the VPRN routing tables dynamically and minimize wasted bandwidth.

The following figure shows a setup with a directly connected IP target.

Figure 11: Directly connected IP target (MD-CLI)

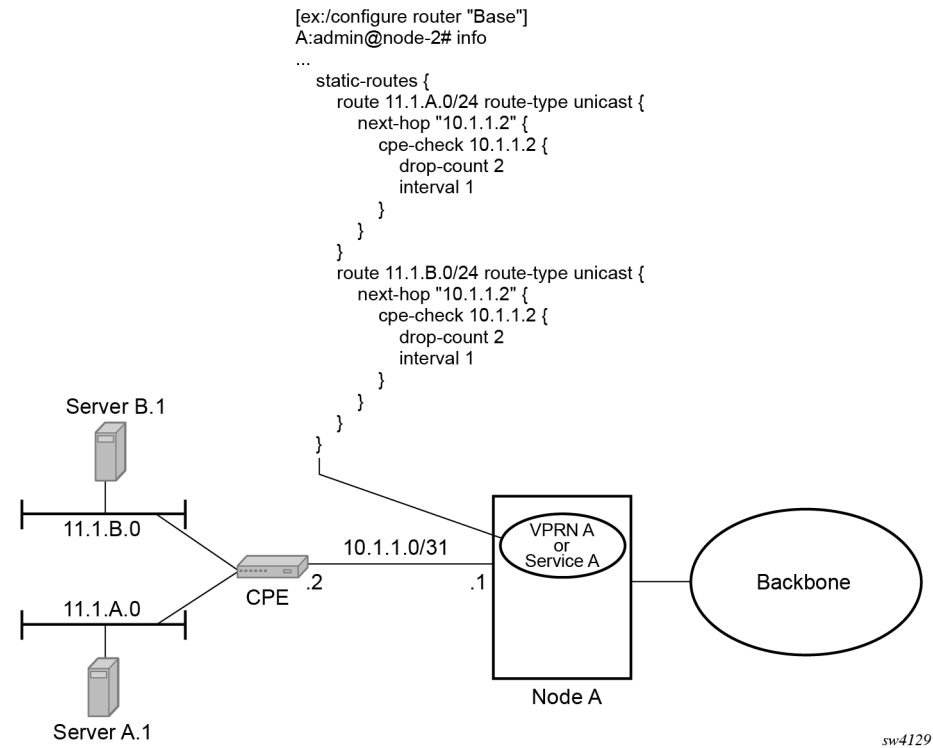
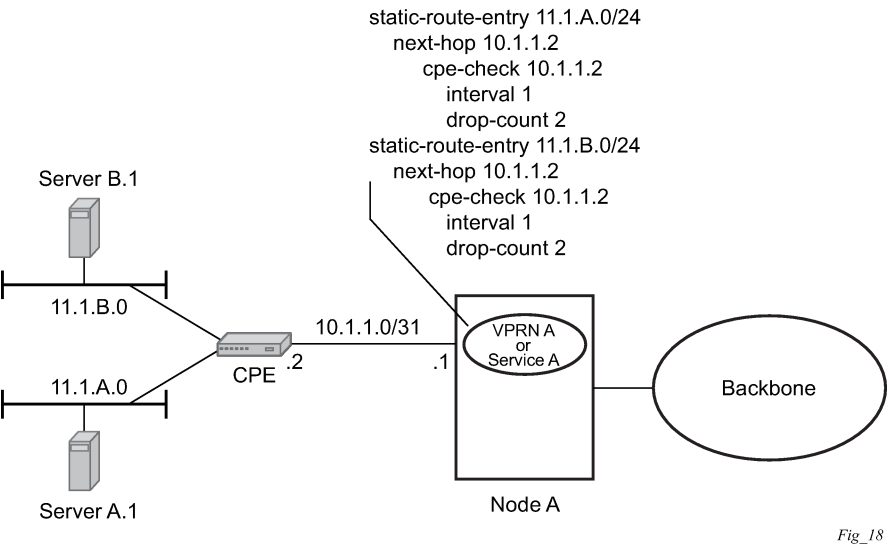


Figure 12: Directly connected IP target (classic CLI)



The following figure shows the setup with multiple hops to an IP target.

Figure 13: Multiple hops to IP target (MD-CLI)

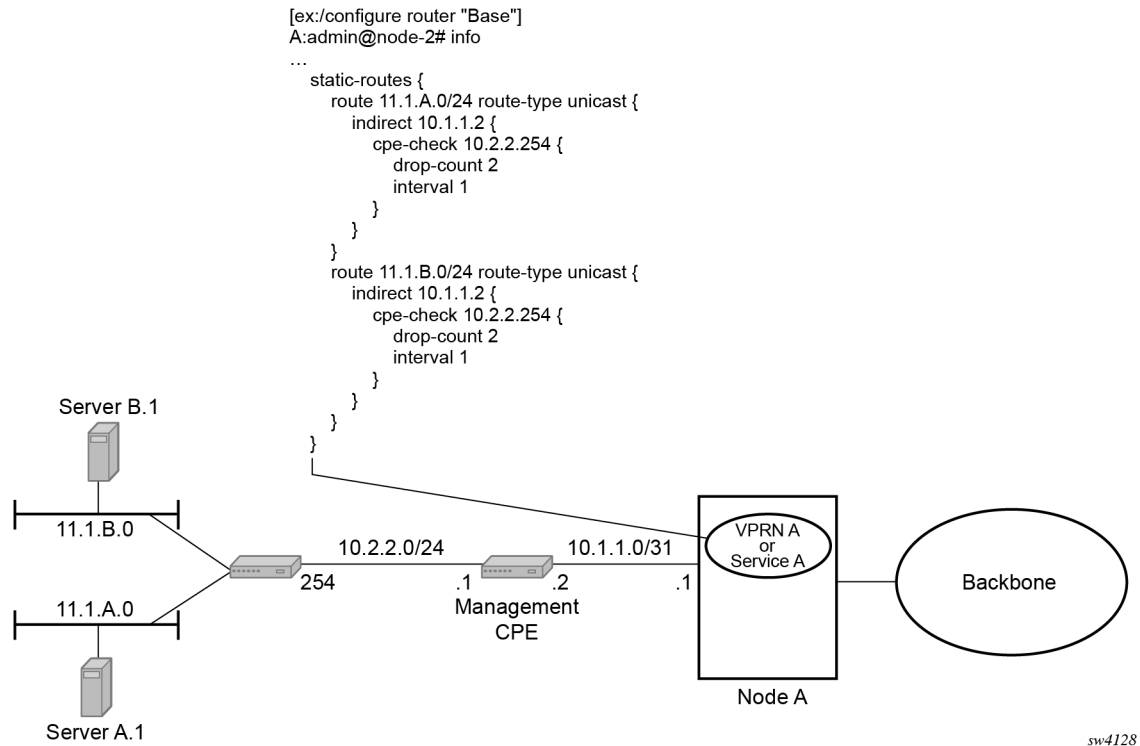
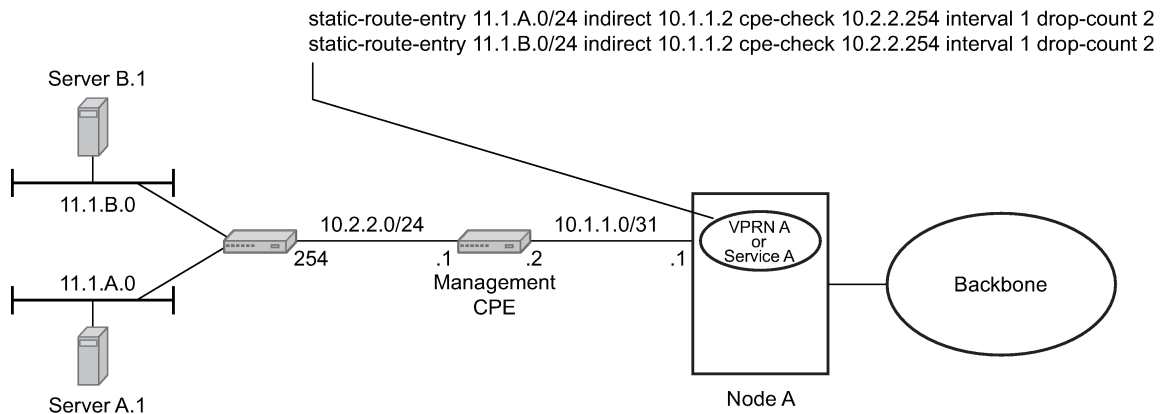


Figure 14: Multiple hops to IP target (classic CLI)



The availability of the far-end static route is monitored through periodic polling. The polling period is configured. If the poll fails a specified number of sequential polls, the static route is marked as inactive. Either ICMP ping or unicast ARP mechanism can be used to test the connectivity. ICMP ping is preferred. If the connectivity check fails and the static route is deactivated, the router continues to send polls and reactivate any routes that are restored.

3.1.6 Constrained route distribution

This section provides information about VPRN Constrained Route Distribution or Route Target Constraint (RTC).

3.1.6.1 Constrained VPN route distribution based on route targets

RTC is a mechanism that allows a router to advertise Route Target (RT) membership information to its BGP peers to indicate interest in receiving only VPN routes tagged with specific RT extended communities. Upon receiving this information, peers restrict the advertised VPN routes to only those requested, minimizing control plane load in terms of protocol traffic and possibly also RIB memory.

The RT membership information is carried in MP-BGP messages, using AFI value 1 and SAFI value of 132. In order for two routers to exchange RT membership Network Layer Reachability Information (NLRI) they must advertise the corresponding AFI/SAFI pair to each other during capability negotiation. By using MP-BGP, RT membership NLRI are propagated, loop-free, within an AS and between ASes using well-known BGP route selection and advertisement rules.

Outbound Route Filtering (ORF) can also be used for RT-based route filtering, but ORF messages have a limited scope of distribution (to direct peers) and, consequently, do not automatically create pruned inter-cluster and inter-AS route distribution trees.

3.1.6.2 Configuring the route target address family

RTC is supported only by the Base router BGP instance. When the **family** command at the BGP router group or neighbor CLI context includes the **route-target** keyword, the RTC capability is negotiated with the associated set of EBGP and IBGP peers.

ORF is mutually exclusive with RTC on a specific BGP session. The CLI does not attempt to block this configuration, but if both capabilities are enabled on a session, the ORF capability is not included in the OPEN message sent to the peer.

3.1.6.3 Originating RTC routes

When the Base router has one or more RTC peers (BGP peers with which the RT Constraint capability has been successfully negotiated), one RTC route is created for each RT extended community imported into a locally-configured L2 VPN or L3 VPN service. Use the following command to configure these imported route targets.

```
configure service vprn
configure service vprn mvprn
```

By default, these RTC routes are automatically advertised to all RTC peers, without the need for an export policy to explicitly "accept" them. Each RTC route has a prefix, a prefix length and path attributes. The prefix value is the concatenation of the origin AS (a 4-byte value representing the 2- or 4-octet AS of the originating router, as configured using the **configure router autonomous-system** command) and 0 or 16-64 bits of a route target extended community encoded in one of the following formats: 2-octet AS specific extended community, IPv4 address specific extended community, or 4-octet AS specific extended community.

Use the following commands to configure a router to send the default RTC route to any RTC peer.

- **MD-CLI**

```
configure router bgp group default-route-target true
configure router bgp group neighbor default-route-target true
```

- **classic CLI**

```
configure router bgp group default-route-target
configure router bgp group neighbor default-route-target
```

The default RTC route is a special type of RTC route that has zero prefix length. Sending the default RTC route to a peer conveys a request to receive all VPN routes (regardless of route target extended community) from that peer. The default RTC route is typically advertised by a route reflector to its clients. The advertisement of the default RTC route to a peer does not suppress other more specific RTC routes from being sent to that peer.

3.1.6.4 Receiving and re-advertising RTC routes

All received RTC routes that are deemed valid are stored in the RIB-IN. An RTC route is considered invalid and treated as withdrawn, if any of the following applies:

- prefix length is 1-31
- prefix length is 33-47
- prefix length is 48-96 and the 16 most-significant bits are not 0x0002, 0x0102 or 0x0202

If multiple RTC routes are received for the same prefix value, standard BGP best path selection procedures are used to determine the best route.



Note: These advertisement rules do not handle hierarchical RR topologies properly. This is a limitation of the current RTC standard.

The best RTC route per prefix is re-advertised to RTC peers based on the following rules:

- The best path for a default RTC route (prefix-length 0, origin AS only with prefix-length 32, or origin AS plus 16 bits of an RT type with prefix-length 48) is never propagated to another peer.
- A PE with only IBGP RTC peers that is neither a route reflector or an ASBR does not re-advertise the best RTC route to any RTC peer because of standard IBGP split horizon rules.
- A route reflector that receives its best RTC route for a prefix from a client peer re-advertises that route (subject to export policies) to all of its client and non-client IBGP peers (including the originator), in accordance with standard RR operation. When the route is re-advertised to client peers, the RR (i) sets the ORIGINATOR_ID to its own router ID and (ii) modifies the NEXT_HOP to be its local address for the sessions (for example, system IP).
- A route reflector that receives its best RTC route for a prefix from a non-client peer re-advertises that route (subject to export policies) to all of its client peers, per standard RR operation. If the RR has a non-best path for the prefix from any of its clients, it advertises the best of the client-advertised paths to all non-client peers.
- An ASBR that is neither a PE nor a route reflector that receives its best RTC route for a prefix from an IBGP peer re-advertises that route (subject to export policies) to its EBGP peers. It modifies the

NEXT_HOP and AS_PATH of the re-advertised route per standard BGP rules. No aggregation of RTC routes is supported.

- An ASBR that is neither a PE nor a route reflector that receives its best RTC route for a prefix from an EBGP peer re-advertises that route (subject to export policies) to its EBGP and IBGP peers. When re-advertised routes are sent to EBGP peers, the ASBR modifies the NEXT_HOP and AS_PATH per standard BGP rules. No aggregation of RTC routes is supported.

3.1.6.5 Using RTC routes

In general (ignoring IBGP-to-IBGP rules, Add-Path, Best-external, and so on), the best VPN route for every prefix/NLRI in the RIB is sent to every peer that supports the VPN address family, but export policies may be used to prevent some prefix/NLRI from being advertised to specific peers. The export policies may be configured statically or created dynamically based on use of ORF or RTC with a peer. ORF and RTC are mutually exclusive on a session.

When RTC is configured on a session that also supports VPN address families using route targets (that is: vpn-ipv4, vpn-ipv6, l2-vpn, mvpn-ipv4, or mvpn-ipv6), the advertisement of the VPN routes is affected as described in the following list:

- When the session comes up, the advertisement of the VPN routes is delayed for a short while to allow RTC routes to be received from the peer.
- After the initial delay, the received RTC routes are analyzed and acted upon. If S1 is the set of routes previously advertised to the peer and S2 is the set of routes that should be advertised based on the most recent received RTC routes the following applies:
 - set of routes in S1 but not in S2 should be withdrawn immediately (subject to MRAI)
 - set of routes in S2 but not in S1 should be advertised immediately (subject to MRAI)
- If a default RTC route is received from a peer P1, the VPN routes that are advertised to P1 is the set that meets the following criteria:
 - routes are eligible for advertisement to P1 per BGP route advertisement rules AND
 - have not been rejected by manually configured export policies AND
 - have not been advertised to the peer



Note: This criterion applies regardless of whether P1 advertised the best route for the default RTC prefix.

In this context, a default RTC route is any of the following:

- route with NLRI length = zero
- route with NLRI value = origin AS and NLRI length = 32
- route with NLRI value = {origin AS+0x0002 | origin AS+0x0102 | origin AS+0x0202} and NLRI length = 48
 - If an RTC route for prefix A (origin-AS = A1, RT = A2/n, n > 48) is received from an IBGP peer I1 in autonomous system A1, the VPN routes that are advertised to I1 is the set that:
 - are eligible for advertisement to I1 per BGP route advertisement rules AND
 - have not been rejected by manually configured export policies AND

- carry at least one route target extended community with value A2 in the n most significant bits AND
- have not been advertised to the peer



Note: This applies regardless of whether I1 advertised the best route for A.

- If the best RTC route for a prefix A (origin-AS = A1, RT = A2/n, $n > 48$) is received from an IBGP peer I1 in autonomous system B, the VPN routes that are advertised to I1 is the set that meets the following criteria:
 - routes are eligible for advertisement to I1 per BGP route advertisement rules AND
 - have not been rejected by manually configured export policies AND
 - carry at least one route target extended community with value A2 in the n most significant bits AND
 - have not been advertised to the peer



Note: This applies only if I1 advertised the best route for A.

- If the best RTC route for a prefix A (origin-AS = A1, RT = A2/n, $n > 48$) is received from an EBGP peer E1, the VPN routes that are advertised to E1 is the set that:
 - are eligible for advertisement to E1 per BGP route advertisement rules AND
 - have not been rejected by manually configured export policies AND
 - carry at least one route target extended community with value A2 in the n most significant bits AND
 - have not been advertised to the peer



Note: This applies only if E1 advertised the best route for A.

3.1.7 BGP fast reroute in a VPRN

BGP fast reroute is a feature that brings together indirection techniques in the forwarding plane and pre-computation of BGP backup paths in the control plane to support fast reroute of BGP traffic around unreachable/failed next-hops. In a VPRN context BGP fast reroute is supported using unlabeled IPv4, unlabeled IPv6, VPN-IPv4, and VPN-IPv6 VPN routes. The supported VPRN scenarios are described in [Table 3: BGP fast reroute scenarios \(VPRN context\)](#).

BGP fast reroute information specific to the Base router BGP context is described in the BGP Fast Reroute section of the *7705 SAR Gen 2 Unicast Routing Protocols Guide*.

Table 3: BGP fast reroute scenarios (VPRN context)

Ingress packet	Primary route	Backup route	Prefix independent convergence
IPv4 (ingress PE)	IPv4 route with next-hop A resolved by an IPv4 route	IPv4 route with next-hop B resolved by an IPv4 route	Yes
IPv4 (ingress PE)	VPN-IPv4 route with next-hop A resolved by a GRE, LDP, RSVP or BGP tunnel	VPN-IPv4 route with next-hop A resolved by a GRE, LDP, RSVP or BGP tunnel	Yes
MPLS (egress PE)	IPv4 route with next-hop A resolved by an IPv4 route	IPv4 route with next-hop B resolved by an IPv4 route	Yes
MPLS (egress PE)	IPv4 route with next-hop A resolved by an IPv4 route	VPN-IPv4 route* with next-hop B resolved by a GRE, LDP, RSVP or BGP tunnel	Yes
IPv6 (ingress PE)	IPv6 route with next-hop A resolved by an IPv6 route	IPv6 route with next-hop B resolved by an IPv6 route	Yes
IPv6 (ingress PE)	VPN-IPv6 route with next-hop A resolved by a GRE, LDP, RSVP or BGP tunnel	VPN-IPv6 route with next-hop B resolved by a GRE, LDP, RSVP or BGP tunnel	Yes
MPLS (egress)	IPv6 route with next-hop A resolved by an IPv6 route	IPv6 route with next-hop B resolved by an IPv6 route	Yes
MPLS (egress)	IPv6 route with next-hop A resolved by an IPv6 route	Yes	VPRN label mode must be VRF. VPRN must export its VPN-IP routes with RD ≠ y. For the best performance the backup next-hop must advertise the same VPRN label value with all routes (per VRF label).

3.1.8 Export of inactive VPRN BGP routes

A BGP route learned from a VPRN BGP peer is exportable as a VPN-IP route, only if it is the best route for the prefix and is installed in the route table of the VPRN. Use the following command to relax this rule and allow the best inactive VPRN BGP route to be exportable as a VPN-IP route, provided that the active installed route for the prefix is an imported VPN-IP route.

```
configure service vprn export-inactive-bgp
```

Use the following command to further relax the preceding rule and allow the best inactive VPRN BGP route (best amongst all routes received from all CEs) to be exportable as a VPN-IP route, regardless of the route type of the active installed route.

```
configure service vprn export-inactive-bgp-enhanced
```

The **export-inactive-bgp** command (or the **export-inactive-bgp-enhanced** command, which is a superset of the functionality) is useful in a scenario where two or more PE routers connect to a multihomed site, and the failure of one of the PE routers or a PE-CE link must be handled by rerouting the traffic over the alternate paths. The traffic failover time in this situation can be reduced if all PE routers have prior knowledge of the potential backup paths and do not have to wait for BGP route advertisements or withdrawals to reprogram their forwarding tables. Achieving this can be challenging with normal BGP procedures. A PE router is not allowed to advertise a BGP route that it has learned from a connected CE device to other PE routers, if that route is not its active route for the destination in the route table. If the multihoming scenario requires all traffic destined for an IP prefix to be carried over a preferred primary path (passing through PE1-CE1, for example), all other PE routers (PE2, PE3, and so on) have that VPN route set as their active route for the destination and are unable to advertise their own routes for the same IP prefix.

For the best inactive BGP route to be exported, it must be matched and accepted by the VRF export policy or there must be an equivalent VRF target configuration.

When a VPN-IP route is advertised because either the **export-inactive-bgp** command or the **export-inactive-bgp-enhanced** command is enabled, the label carried in the route is either a per next-hop label corresponding to the next-hop IP address of the CE-BGP route, or a per-prefix label. This helps to avoid packet looping issues caused by unsynchronized IP FIBs.

When a PE router that advertised an inactive VPRN BGP route for an IP prefix receives a withdrawal for the VPN-IP route that was the active primary route, the inactive backup path may be promoted to the primary path; that is, the CE-BGP route may become the active route for the destination. In this case, the PE router is required to readvertise the VPN-IP route with a per-VRF label, if that is the behavior specified by the default allocation policy, and there is no label-per-prefix policy override. In the time it takes for the new VPN-IP route to reach all ingress routers and for the ingress routers to update their forwarding tables, traffic continues to be received with the old per next-hop label. The egress PE drops the in-flight traffic, unless the following command is used to configure label retention.

```
configure router mpls-labels bgp-labels-hold-timer
```

The **bgp-labels-hold-timer** command configures a delay (in seconds) between the withdrawal of a VPN-IP route with a per next-hop label and the deletion of the corresponding label forwarding entry in the IOM. The value of the **bgp-labels-hold-timer** command must be large enough to account for the propagation delay of the route withdrawal to all ingress routers.

3.2 VPRN features

This section describes the VPRN features and any special capabilities or considerations as they relate to VPRN services.

3.2.1 IP interfaces

VPRN customer IP interfaces can be configured with most of the same options found on the base IP interfaces.

The supported advanced configuration options are as follows:

- Cflowd
- VRRP
- secondary IP addresses
- ICMP

NTP broadcast receipt, one of the base IP interface configuration options, is not supported on VPRN IP interfaces.

3.2.1.1 Displaying QoS information associated with routes

Use the commands in the following contexts to show the forwarding class and priority associated with the displayed routes.

```
show router route-table
show router fib
show router bgp routes
show router rip database
show router static-route
```

Use the following command to show an additional line per route entry that displays the forwarding class and priority of the route. When the **qos** command option is specified, the output includes an additional line per route entry that displays the forwarding class and priority of the route. If a route has no forwarding class and priority information, the third line is blank.

```
show router route-table 10.1.5.0/24 qos
```

Output example

```
=====
Route Table (Router: Base)
=====
Dest Prefix                                Type   Proto   Age      Metric   Pref
  Next Hop[Interface Name]
    QoS
-----
10.1.5.0/24                               Remote BGP     15h32m52s  0
  PE1_to_PE2
  h1, high
-----
No. of Routes: 1
=====
```

3.2.1.2 Object grouping and state monitoring

This feature introduces a generic operational group object which associates different service endpoints (pseudowires and SAPs) located in the same or in different service instances. The operational group status is derived from the status of the individual components using specific rules specific to the application using the concept. A number of other service entities, the monitoring objects, can be configured to monitor the operational group status and to perform specific actions as a result of status transitions. For example, if the operational group goes down, the monitoring objects are brought down.

3.2.1.3 VPRN IP interface applicability

About this task

This concept is used by an IPv4 VPRN interface to affect the operational state of the IP interface monitoring the operational group. Individual SAP and spoke SDPs are supported as monitoring objects.

The following rules apply:

- An object can only belong to one group at a time.
- An object that is part of a group cannot monitor the status of a group.
- An object that monitors the status of a group cannot be part of a group.
- An operational group may contain any combination of member types: SAP or spoke SDPs.
- An operational group may contain members from different VPLS service instances.
- Objects from different services may monitor the operational group.

Procedure

Step 1. Identify a set of objects whose forwarding state should be considered as a whole group then group them under an operational group using the **oper-group** command.

Step 2. Associate the IP interface to the operational group using the **monitor-oper-group** command. The status of the operational group is dictated by the status of one or more members according to the following rules:

- The operational group goes down if all the objects in the operational group go down. The operational group comes up if at least one of the components is up.
- An object in the group is considered down if it is not forwarding traffic in at least one direction. That could be because the operational state is down or the direction is blocked through some validation mechanism.
- If a group is configured but no members are specified yet, its status is considered up.
- As soon as the first object is configured the status of the operational group is dictated by the status of the provisioned members.

The following configuration shows the operational group g1, the VPLS SAP that is mapped to it and the IP interfaces in IES service 2001 monitoring the operational group g1. This example uses an R-VPLS context.

Operational group g1 has a single SAP (1/1/1:2001) mapped to it and the IP interfaces in the IES service 2001 derive their state from the state of operational group g1.

Example**MD-CLI**

In the MD-CLI, the VPLS instance includes the name v1. The IES interface links to the VPLS using the **vpls** command.

```
[ex:/configure service]
A:admin@node-2# info
  oper-group "g1" {
  }
  ies "2001" {
    customer "1"
    interface "i2001" {
      monitor-oper-group "g1"
      vpls "v1" {
      }
      ipv4 {
        primary {
          address 192.168.1.1
          prefix-length 24
        }
      }
    }
  }
}
vpls "v1" {
  admin-state enable
  service-id 1
  customer "1"
  routed-vpls {
  }
  stp {
    admin-state disable
  }
  sap 1/1/1:2001 {
    oper-group "g1"
    eth-cfm {
      mep md-admin-name "1" ma-admin-name "1" mep-id 1 {
      }
    }
  }
  sap 1/1/2:2001 {
    admin-state enable
  }
  sap 1/1/3:2001 {
    admin-state enable
  }
}
```

Example**Classic CLI**

In the classic CLI, the VPLS instance includes the **allow-ip-int-bind** and the name v1. The IES interface links to the VPLS using the **vpls** command.

```
A:node-2>config>service# info
-----
  oper-group "g1" create
  exit
  vpls 1 name "v1" customer 1 create
    allow-ip-int-bind
    exit
    stp
      shutdown
```

```

exit
sap 1/1/1:2001 create
oper-group "g1"
eth-cfm
    mep 1 domain 1 association 1 direction down
    ccm-enable
no shutdown
exit
sap 1/1/2:2001 create
no shutdown
exit
sap 1/1/3:2001 create
no shutdown
exit
no shutdown
exit
ies 2001 name "2001" customer 1 create
shutdown
interface "i2001" create
    monitor-oper-group "g1"
    address 192.168.1.1/24
    vpls "v1"
exit
exit
exit

```

3.2.2 SAPs

3.2.2.1 SAP encapsulations

The following SAP encapsulations are supported on the VPRN service:

- Ethernet null
- Ethernet dot1q
- QinQ
- LAG
- Tunnel (IPsec or GRE)

3.2.2.2 TCP MSS adjustment filter on VPRN SAP interfaces

About this task

VSR supports a configurable filter that adjusts the maximum segment size (MSS) of TCP packets marked with a SYN flag that traverse VPRN SAP interfaces. The MSS adjustment filter prevents upstream and downstream TCP packets from being fragmented.

MSS adjustment is performed by the virtualized integrated BB ISA MDA when an IP filter is enabled with the **configure filter ip-filter entry action tcp-mss-adjust** command. The filter can be applied to a VPRN SAP interface in the ingress, egress, or both directions. Both IPv4 and IPv6 filters are supported. See the *7450 ESS and 7750 SR Multiservice ISA and ESA Advanced Configuration Guide for Classic CLI* for information about the virtualized BB ISA MDA.

Perform the following steps to configure a a TCP MSS adjustment filter on a VPRN SAP interface.

The CLI output in the following steps is an example of the creation of a NAT group on the virtualized integrated BB ISA MDA in slot 1/6.

Procedure

Step 1. Configure the BB ISA MDA type.

Example

Configure BB ISA MDA type for slot 1/6 (MD-CLI)

```
[ex:/configure card 1]
A:admin@node-2# info
card-type iom-v
mda 6 {
    mda-type isa-bb-v
}
```

Example

Configure BB ISA MDA type for slot 1/6 (classic CLI)

```
A:node-2>config# card 1
A:node-2>config>card# info
-----
card-type iom-v
mda 6
    mda-type isa-bb-v
    no shutdown
exit
no shutdown
-----
```

Step 2. Configure the NAT group.

Example

Configure NAT group (MD-CLI)

```
[ex:/configure isa]
A:admin@node-2# info
nat-group 1 {
    admin-state enable
    redundancy {
        active-mda-limit 1
    }
mda 1/6 { }
```

Example

Configure NAT group (classic CLI)

```
A:node-2>config>isa# info
-----
nat-group 1 create
active-mda-limit 1
mda 1/6
no shutdown
exit
-----
```

- Step 3.** Associate the NAT group with a routing instance and configure the MSS value as shown in the following example.

Example

Associate the NAT group with a routing instance and configure the MSS value (MD-CLI)

```
[ex:/configure service vprn "2"]
A:admin@node-2# info
  customer "1"
  mss-adjust {
    nat-group 1
    segment-size 1352
  }
```

Example

Associate the NAT group with a routing instance and configure the MSS value (classic CLI)

```
A:node-2>config>service>vprn# info
-----
      shutdown
      mss-adjust-group 1 segment-size 1352
-----
```

- Step 4.** Create ingress or egress IP filters that perform TCP MSS adjustment.

Example

Configure IPv4 filters and IPv6 filters for TCP MSS adjustment at ingress and egress (MD-CLI)

```
[ex:/configure filter]
A:admin@node-2# info
  ip-filter "1" {
    description "Ingress"
    default-action accept
    entry 1 {
      match {
        protocol tcp
        tcp-flags {
          syn true
        }
      }
      action {
        tcp-mss-adjust
      }
    }
  }
  ip-filter "2" {
    description "Egress"
    default-action accept
    entry 1 {
      egress-pbr true
      match {
        protocol tcp
        tcp-flags {
          syn true
        }
      }
      action {
        tcp-mss-adjust
      }
    }
  }
```

```

    }
  }
}
ipv6-filter "four" {
  description "Egress"
  default-action accept
  filter-id 2
  entry 1 {
    egress-pbr true
    match {
      next-header tcp
      tcp-flags {
        syn true
      }
    }
    action {
      tcp-mss-adjust
    }
  }
}
ipv6-filter "three" {
  description "Ingress"
  default-action accept
  filter-id 1
  entry 1 {
    match {
      next-header tcp
      tcp-flags {
        syn true
      }
    }
    action {
      tcp-mss-adjust
    }
  }
}
}

```

Example

Configure IPv4 filters and IPv6 filters for TCP MSS adjustment at ingress and egress (classic CLI)

```

A:node-2>config>filter# info
-----
ip-filter 1 name "1" create
  default-action forward
  description "Ingress"
  entry 1 create
    match protocol tcp
    tcp-syn true
  exit
  action
    tcp-mss-adjust
  exit
exit
ip-filter 2 name "2" create
  default-action forward
  description "Egress"
  entry 1 create
    match protocol tcp
    tcp-syn true
  exit
  action

```

```

        tcp-mss-adjust
    exit
    egress-pbr default-load-balancing
exit
exit
ipv6-filter 1 name "three" create
    default-action forward
    description "Ingress"
    entry 1 create
        match next-header tcp
        tcp-syn true
    exit
    action
        tcp-mss-adjust
    exit
exit
exit
ipv6-filter 2 name "four" create
    default-action forward
    description "Egress"
    entry 1 create
        match next-header tcp
        tcp-syn true
    exit
    action
        tcp-mss-adjust
    exit
    egress-pbr default-load-balancing
exit
exit
-----

```

Step 5. Apply the filters that perform TCP MSS adjustment to the VPRN SAP interface. The filters can be applied in the ingress, egress, or both directions. In the following example, the filters are applied in both directions.

Example

Apply the filters for TCP MSS adjustment to the VPRN SAP interface (MD-CLI)

```

[ex:/configure service]
A:admin@node-2# info
    vprn "2" {
        customer "1"
        interface "int1_vprn1" {
            ipv4 {
                primary {
                    address 10.10.1.1
                    prefix-length 24
                }
            }
            sap 1/2/3 {
                ingress {
                    filter ip 1
                }
                egress {
                    filter ip 2
                }
            }
        }
    }
    vprn "3" {
        customer "1"
        interface "int1_vprn3" {

```

```

        ipv6 {
            address 10:1::1 {
                prefix-length 32
            }
            neighbor-discovery {
                static-neighbor 10:1::2 {
                    mac-address 00:02:01:00:00:01
                }
            }
        }
    }
    sap 1/2/3 {
        ingress {
            filter ipv6 3
        }
        egress {
            filter ipv6 4
        }
    }
}

```

Example**Apply the filters for TCP MSS adjustment to the VPRN SAP interface (classic CLI)**

```

A:node-2>config>service# info
-----
vprn 2 name "2" customer 1 create
shutdown
interface "intl_vprn1" create
address 10.10.1.1/24
sap 1/2/3 create
    ingress
        filter ip 1
    exit
    egress
        filter ip 2
    exit
exit
exit
exit
vprn 3 name "3" customer 1 create
shutdown
interface "intl_vprn3" create
    ipv6
        address 10:1::1/32
        neighbor 10:1::2 00:02:01:00:00:01
    exit
    sap 1/2/3 create
        ingress
            filter ipv6 3
        exit
        egress
            filter ipv6 4
        exit
    exit
exit
exit
exit

```

3.2.3 QoS policies

When applied to a VPRN SAP, service ingress QoS policies only create the unicast queues defined in the policy if PIM is not configured on the associated IP interface; if PIM is configured, the multipoint queues are applied as well.

With VPRN services, service egress QoS policies function as with other services where the class-based queues are created as defined in the policy.

Both Layer 2 and Layer 3 criteria can be used in the QoS policies for traffic classification in an VPRN.

3.2.4 Filter policies

Ingress and egress IPv4 and IPv6 filter policies can be applied to VPRN SAPs.

3.2.5 DSCP marking

Specific DSCP, forwarding class, and Dot1P command options can be specified to be used by every protocol packet generated by the VPRN. This enables prioritization or de-prioritization of every protocol (as required). The markings effect a change in behavior on ingress when queuing. For example, if OSPF is not enabled, then traffic can be de-prioritized to best effort (be) DSCP. This change de-prioritizes OSPF traffic to the CPU complex.

DSCP marking for internally generated control and management traffic by marking the DSCP value should be used for the specific application. This can be configured per routing instance. For example, OSPF packets can carry a different DSCP marking for the base instance and then for a VPRN service. ISIS and ARP traffic is not an IP-generated traffic type and is not configurable. See [Table 4: DSCP/FC marking](#).

When an application is configured to use a specified DSCP value then the MPLS EXP, Dot1P bits are marked in accordance with the network or access egress policy as it applies to the logical interface the packet is egressing.

The DSCP value can be set per application. This setting is forwarded to the egress line card. The egress line card does not alter the coded DSCP value and marks the LSP-EXP and IEEE 802.1p (Dot1P) bits according to the appropriate network or access QoS policy.

Table 4: DSCP/FC marking

Protocol	IPv4	IPv6	DSCP marking	Dot1P marking	Default FC
ARP	—	—	—	Yes	NC
BGP	Yes	Yes	Yes	Yes	NC
BFD	Yes	—	Yes	Yes	NC
RIP	Yes	Yes	Yes	Yes	NC
PIM (SSM)	Yes	Yes	Yes	Yes	NC
OSPF	Yes	Yes	Yes	Yes	NC

Protocol	IPv4	IPv6	DSCP marking	Dot1P marking	Default FC
SMTP	Yes	—	—	—	AF
IGMP/MLD	Yes	Yes	Yes	Yes	AF
Telnet	Yes	Yes	Yes	Yes	AF
TFTP	Yes	—	Yes	Yes	AF
FTP	Yes	—	—	—	AF
SSH (SCP)	Yes	Yes	Yes	Yes	AF
SNMP (get, set, and so on)	Yes	Yes	Yes	Yes	AF
SNMP trap/log	Yes	Yes	Yes	Yes	AF
syslog	Yes	Yes	Yes	Yes	AF
OAM ping	Yes	Yes	Yes	Yes	AF
ICMP ping	Yes	Yes	Yes	Yes	AF
Traceroute	Yes	Yes	Yes	Yes	AF
TACPLUS	Yes	Yes	Yes	Yes	AF
DNS	Yes	Yes	Yes	Yes	AF
SNTP/NTP	Yes	—	—	—	AF
RADIUS	Yes	—	—	—	AF
Cflowd	Yes	—	—	—	AF
DHCP	Yes	Yes	Yes	Yes	AF
Bootp	Yes	—	—	—	AF
IPv6 Neighbor Discovery	Yes	—	—	—	NC

3.2.5.1 Default DSCP mapping table

```

DSCP NameDSCP ValueDSCP ValueDSCP ValueLabel
Decimal Hexadecimal Binary
=====
Default 00x00 0b0000000be
nc1 48 0x30 0b110000h1
nc2 56 0x38 0b111000nc
ef 46 0x2e 0b101110ef
af11100x0a0b001010assured

```

```
af12120x0c0b001100assured
af13140x0e0b001110assured
af21 18 0x12 0b01001011
af22 20 0x14 0b01010011
af23220x160b01011011
af31 26 0x1a 0b01101011
af32 28 0x1c 0b01110011
af33 30 0x1d 0b01111011
af41 34 0x22 0b100010h2
af42 36 0x24 0b100100h2
af43 38 0x26 0b100110h2

default*0
```

*The default forwarding class mapping is used for all DSCP names or values for which there is no explicit forwarding class mapping.

3.2.6 Configuration of TTL propagation for VPRN routes

This feature allows the separate configuration of TTL propagation for in transit and CPM generated IP packets, at the ingress LER within a VPRN service context. Use the following commands to configure TTL propagation.

```
configure router ttl-propagate vprn-local
configure router ttl-propagate vprn-transit
```

You can enable TTL propagation behavior separately as follows:

- for locally generated packets by CPM (using the **vprn-local** command)
- for user and control packets in transit at the node (using the **vprn-transit** command)

The following command options can be specified:

- **all** – enables TTL propagation from the IP header into all labels in the stack, for VPN-IPv4 and VPN-IPv6 packets forwarded in the context of all VPRN services in the system.
- **vc-only** – reverts to the default behavior by which the IP TTL is propagated into the VC label but not to the transport labels in the stack. You can explicitly set the default behavior by configuring the **vc-only** value.
- **none** – disables the propagation of the IP TTL to all labels in the stack, including the VC label. This is needed for a transparent operation of UDP traceroute in VPRN Inter-AS Option B such that the ingress and egress ASBR nodes are not traced.

In the classic CLI, the **vprn-local** command does not use a **no** version.

In the MD-CLI, use the **delete** command to remove the configuration.

Use the following commands to override the global configuration within each VPRN instance.

```
configure service vprn ttl-propagate local
configure service vprn ttl-propagate transit
```

The default behavior for a VPRN instance is to inherit the global configuration for the same command. You can explicitly set the default behavior by configuring the **inherit** value.

In the classic CLI, the **local** and **transit** commands do not have **no** versions.

In the MD-CLI, use the **delete** command to remove the configuration.

The commands do not apply when the VPRN packet is forwarded over GRE transport tunnel.

If a packet is received in a VPRN context and a lookup is done in the Global Routing Table (GRT), when leaking to GRT is enabled, for example, the behavior of the TTL propagation is governed by the LSP shortcut configuration as follows:

- when the matching route is an RSVP LSP shortcut

```
configure router mpls shortcut-transit-ttl-propagate
```

- when the matching route is an LDP LSP shortcut

```
configure router ldp shortcut-transit-ttl-propagate
```

When the matching route is a RFC 8277 label route or a 6PE route, It is governed by the BGP label route configuration.

When a packet is received on one VPRN instance and is redirected using Policy Based Routing (PBR) to be forwarded in another VPRN instance, the TTL propagation is governed by the configuration of the outgoing VPRN instance.

Packets that are forwarded in different contexts can use different TTL propagation over the same BGP tunnel, depending on the TTL configuration of each context. An example of this may be VPRN using a BGP tunnel and an IPv4 packet forwarded over a BGP label route of the same prefix as the tunnel.

3.2.7 CE to PE routing protocols

The 7705 SAR Gen 2 supports the following PE to CE routing protocols:

- BGP
- Static
- RIP
- OSPF

3.2.7.1 PE to PE tunneling mechanisms

The 7705 SAR Gen 2 supports multiple mechanisms to provide transport tunnels for the forwarding of traffic between PE routers within the 2547bis network.

The 7705 SAR Gen 2 VPRN implementation supports the use of:

- RSVP-TE protocol to create tunnel LSPs between PE routers
- LDP protocol to create tunnel LSP's between PE routers
- GRE tunnels between PE routers

These transport tunnel mechanisms provide the flexibility of using dynamically created LSPs where the service tunnels are automatically bound (the autobind feature) and the ability to provide specific VPN services with their own transport tunnels by explicitly binding SDPs if needed. When the autobind is used, all services traverse the same LSPs and do not allow alternate tunneling mechanisms (like GRE) or the ability to craft sets of LSPs with bandwidth reservations for specific customers as is available with explicit SDPs for the service.

3.2.7.2 Per VRF route limiting

The 7705 SAR Gen 2 allows setting the maximum number of routes that can be accepted in the VRF for a VPRN service. There are options to specify a percentage threshold at which to generate an event that the VRF table is near full and an option to disable additional route learning when full or only generate an event.

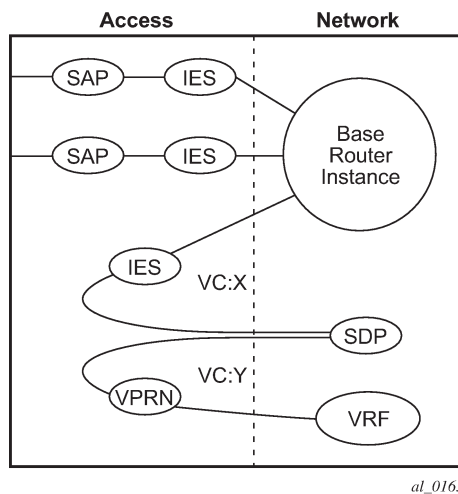
3.2.8 Spoke SDPs

Distributed services use service destination points (SDPs) to direct traffic to another router via service tunnels. SDPs are created on each participating router and then bound to a specific service. SDP can be created as either GRE or MPLS. See the *7705 SAR Gen 2 Services Overview Guide* for information about configuring SDPs.

This feature provides the ability to cross-connect traffic entering on a spoke SDP, used for Layer 2 services (VLLs or VPLS), on to an IES or VPRN service. From a logical point of view, the spoke SDP entering on a network port is cross-connected to the Layer 3 service as if it entered by a service SAP. The main exception to this is traffic entering the Layer 3 service by a spoke SDP is handled with network QoS policies and not access QoS policies.

Figure 15: SDP-ID and VC label service identifiers depicts traffic terminating on a specific IES or VPRN service that is identified by the SDP ID and VC label present in the service packet.

Figure 15: SDP-ID and VC label service identifiers



See "VCCV BFD support for VLL, Spoke SDP Termination on IES and VPRN, and VPLS Services" in the *7705 SAR Gen 2 Layer 2 Services and EVPN Guide* for information about using VCCV BFD in spoke-SDP termination.



Note: Spoke-SDP termination of Ipipe VLLs on VPRN is not supported in System Profile B. Use the following command to determine if Ipipes are currently bound to an VPRN interface, before configuring System Profile B.

```
show router ldp bindings services
```

3.2.8.1 T-LDP status signaling for spoke-SDPs terminating on IES/VPRN

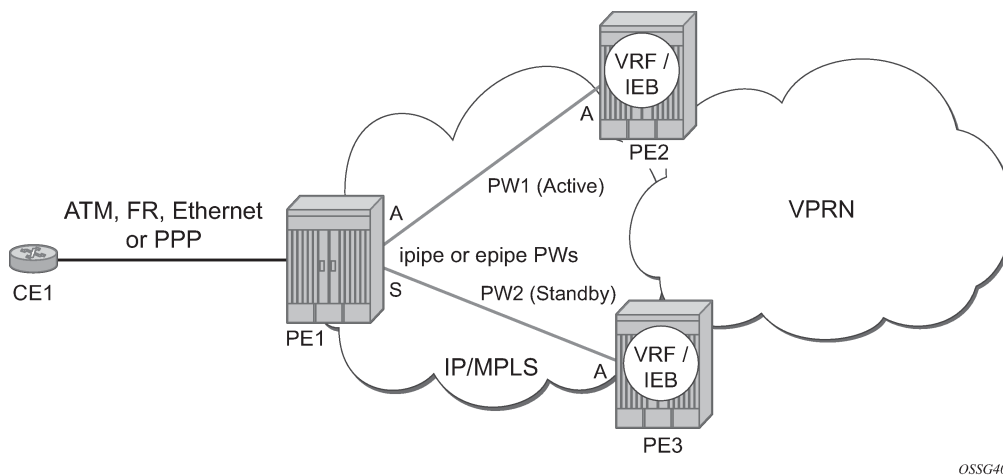
T-LDP status signaling and PW active/standby signaling capabilities are supported on Ipipe and Epipe spoke SDPs.

Spoke SDP termination on an IES or VPRN provides the ability to cross-connect traffic entering on a spoke SDP, used for Layer 2 services (VLLs or VPLS), on to an IES or VPRN service. From a logical point of view the spoke SDP entering on a network port is cross-connected to the Layer 3 service as if it had entered using a service SAP. The main exception to this is traffic entering the Layer 3 service using a spoke SDP is handled with network QoS policies instead of access QoS policies.

When a SAP down or SDP binding down status message is received by the PE in which the Ipipe or Ethernet Spoke-SDP is terminated on an IES or VPRN interface, the interface is brought down and all associated routes are withdrawn in a similar way when the Spoke-SDP goes down locally. The same actions are taken when the standby T-LDP status message is received by the IES/VPRN PE.

This feature can be used to provide redundant connectivity to a VPRN or IES from a PE providing a VLL service, as shown in [Figure 16: Active/standby VRF using resilient Layer 2 circuits](#).

Figure 16: Active/standby VRF using resilient Layer 2 circuits



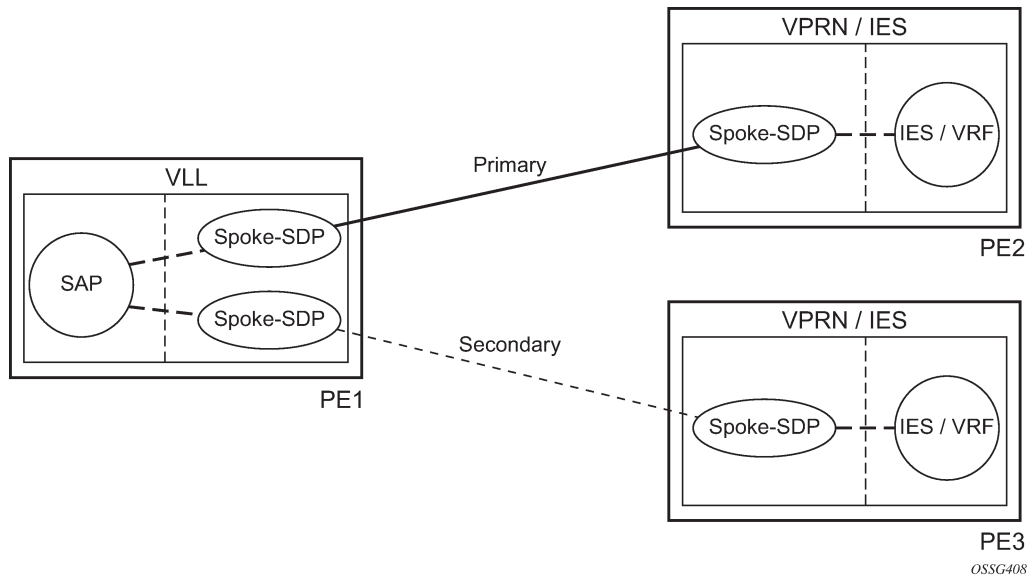
3.2.8.2 Spoke SDP redundancy into IES/VPRN

This feature can be used to provide redundant connectivity to a VPRN or IES from a PE providing a VLL service, as shown in [Figure 16: Active/standby VRF using resilient Layer 2 circuits](#), using either Epipe or Ipipe spoke-SDPs.

In [Figure 16: Active/standby VRF using resilient Layer 2 circuits](#), PE1 terminates two spoke SDPs that are bound to one SAP connected to CE1. PE1 chooses to forward traffic on one of the spoke SDPs (the active spoke-SDP), while blocking traffic on the other spoke SDP (the standby spoke SDP) in the transmit direction. PE2 and PE3 take any spoke SDPs for which PW forwarding standby has been signaled by PE1 to an operationally down state.

[Figure 17: Spoke SDP redundancy model](#) illustrates the model for spoke SDP redundancy into a VPRN or IES.

Figure 17: Spoke SDP redundancy model



3.2.8.3 Weighted ECMP for spoke-SDPs terminating on IES/VPRN and R-VPLS interfaces

ECMP and weighted ECMP into RSVP-TE and SR-TE LSPs is supported for Ipipe and Epipe spoke SDPs terminating on IP interfaces in an IES or VPRN, or for spoke SDP termination on a routed VPLS. It is also supported for SDPs using LDP over RSVP tunnels. The following example shows the configuration of weighted ECMP under the SDP used by the service.

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2# info
...
sdp 1 {
  delivery-type mpls
  weighted-ecmp true
}
```

Example: classic CLI

```
A:node-2>config>service# info
-----
sdp 1 mpls create
shutdown
weighted-ecmp
exit
```

When a service uses a provisioned SDP on which weighted ECMP is configured, a path is selected based on the configured hash. Paths are then load balanced across LSPs used by an SDP according to normalized LSP load balancing weights. If one or more LSPs in the ECMP set to a specific next hop has no **load-balancing-weight** value configured, regular ECMP spraying is used.

3.2.9 IP-VPNs

3.2.9.1 Using OSPF in IP-VPNs

Using OSPF as a CE to PE routing protocol allows OSPF that is currently running as the IGP routing protocol to migrate to an IP-VPN backbone without changing the IGP routing protocol, introducing BGP as the CE-PE or relying on static routes for the distribution of routes into the service providers IP-VPN. The following features are supported:

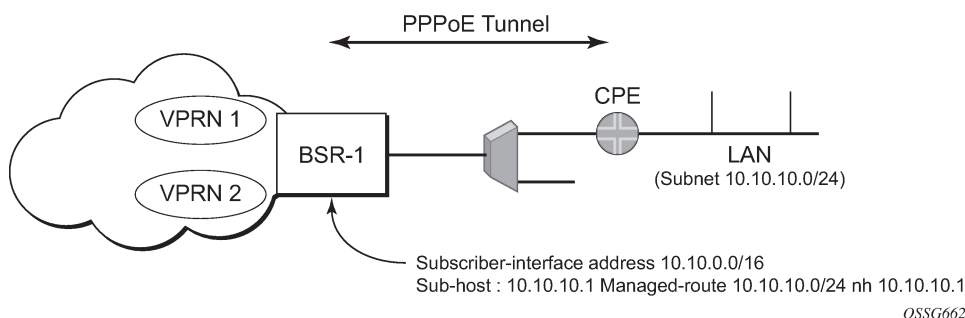
- Advertisement/redistribution of BGP-VPN routes as summary (type 3) LSAs flooded to CE neighbors of the VPRN OSPF instance. This occurs if the OSPF route type (in the OSPF route type BGP extended community attribute carried with the VPN route) is not external (or NSSA) and the locally configured domain-id matches the domain-id carried in the OSPF domain ID BGP extended community attribute carried with the VPN route.
- OSPF sham links; a sham link is a logical PE-to-PE unnumbered point-to-point interface that essentially rides over the PE-to-PE transport tunnel. A sham link can be associated with any area and can therefore appear as an intra-area link to CE routers attached to different PEs in the VPN.

3.2.10 IPCP subnet negotiation

This feature enables negotiation between Broadband Network Gateway (BNG) and customer premises equipment (CPE) so that CPE is allocated to both the IP address and associated subnet.

Some CPEs use the network up-link in PPPoE mode and perform the DHCP server function for all ports on the LAN side. Instead of wasting 1 subnet for p2p uplink, CPEs use the allocated subnet for the LAN portion as shown in [Figure 18: CPEs network up-link mode](#).

Figure 18: CPEs network up-link mode



From a BNG perspective, the specific PPPoE host is allocated a subnet (instead of /32) by RADIUS, external DHCP server, or local user database. Locally, the host is associated with a managed route. This managed route is a subset of the subscriber-interface subnet on the 7705 SAR Gen 2, and the subscriber host IP address is from the managed-route range. The negotiation between BNG and CPE allows CPE to be allocated both an IP address and associated subnet.

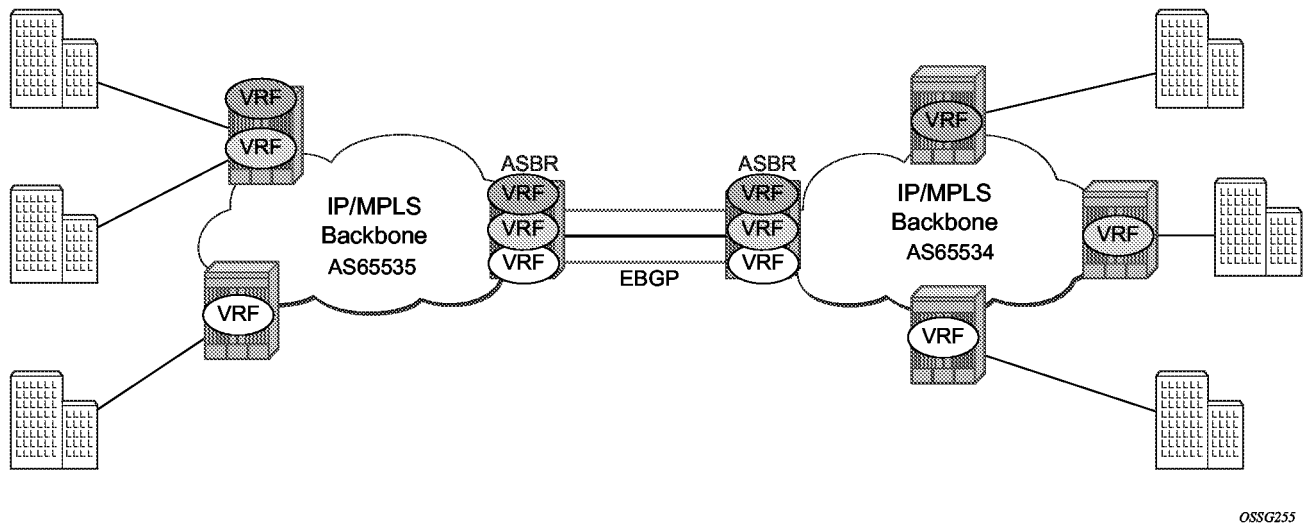
3.2.11 Inter-AS VPRNs

Inter-AS IP-VPN services enable service providers to expand beyond a single Autonomous System (AS) and IP-VPN services to cross the AS boundaries of multiple providers. The inter-AS IP-VPN options are described in RFC 4364, *BGP/MPLS IP Virtual Private Networks (VPNs)*.

3.2.11.1 Inter-AS option-A

Inter-AS option-A is inherent in any implementation. As shown in the following figure, option-A uses a back-to-back connection between separate VPRN instances in each AS. As a result, each VPRN instance views the inter-AS connection as an external interface to a remote VPRN customer site. The back-to-back VRF connections between the ASBR nodes require individual sub-interfaces, one for each VRF.

Figure 19: Inter-AS option-A: VRF-to-VRF model



OSSG255

3.2.11.2 Inter-AS Option B

The recursive opaque type used for Inter-AS Option B is the Recursive Opaque (VPN Type), shown as opaque type 8 in [Table 8: Recursive opaque types](#).

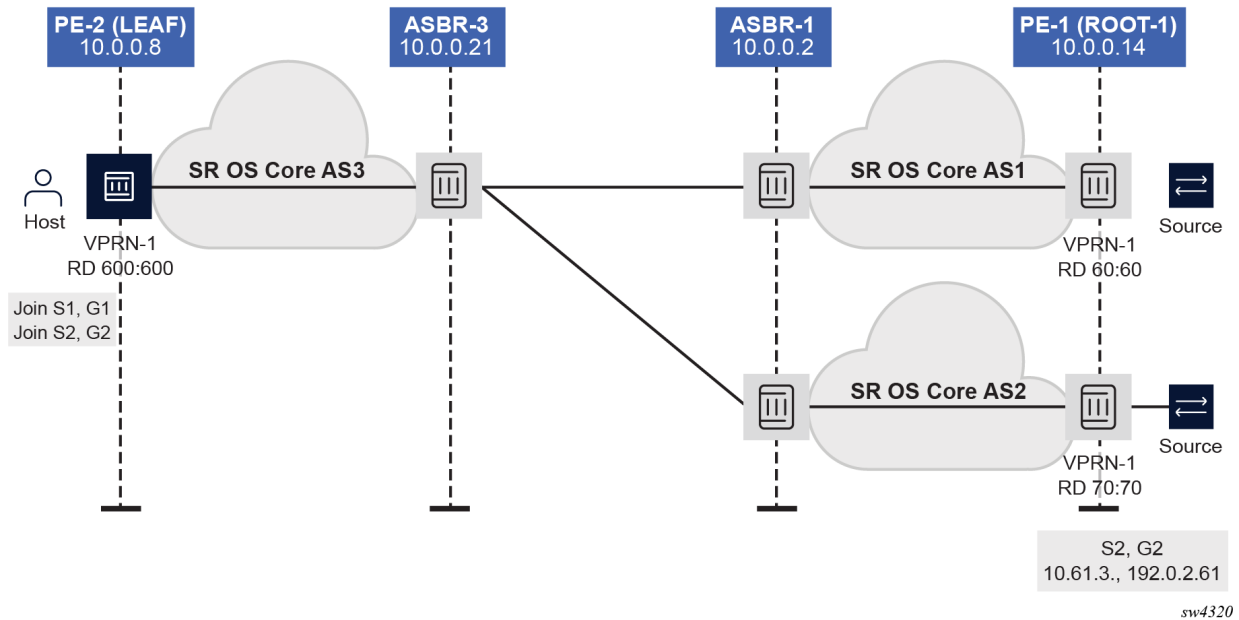
Inter-AS option B support for NG-MVPN

Inter-AS Option B requires additional processing on ASBR routers and recursive FEC encoding than that of Inter-AS Option A. Because BGP adjacency is not e2e, ASBRs must cache and use a PMSI route to build the tree. For that, mLDP recursive FEC must carry RD information—therefore, VPN recursive FEC is required (opaque type 8).

In Inter-AS Option B, the PEs in two different ASs do not have their system IP address in the RTM. As such, for NG-MVPN, a recursive opaque value in mLDP FEC is required to signal the LSP to the first ASBR in the local AS path.

Because the system IPs of the peer PEs (Root-1 and Root-2) are not installed on the local PE (leaf), it is possible to have two PEs in different ASs with same system IP address, as shown in [Figure 20: Identical system IP on multiple PEs \(Option B\)](#). However, SR OS does not support this topology. The system IP address of all nodes (root or leaf) in different ASs must be unique.

Figure 20: Identical system IP on multiple PEs (Option B)



For Inter-AS Option B and NG-MVPN, SR OS as a leaf does not support multiple roots in multiple ASs with the same system IP and different RDs; however, the first root that is advertised to an SR OS leaf is used by PIM to generate an MLDP tunnel to this actual root. Any dynamic behavior after this point, such as removal of the root and its replacement by a second root in a different AS, is not supported and the SR OS behavior is nondeterministic.

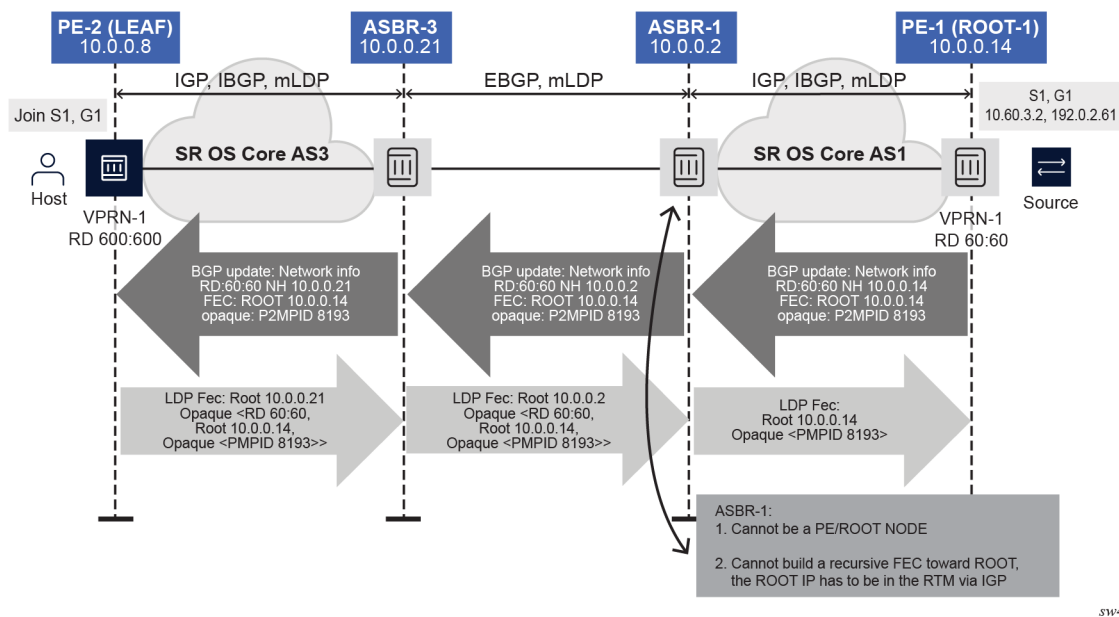
I-PMSI and S-PMSI establishment

I-PMSI and S-PMSI functionality follows RFC 6513 section 8.1.1 and RFC 6512 sections 3.1 and 3.2.1. For routing, the same rules as for GRT d-MLDP use case apply, but the VRR Route Import External community now encodes the VRF instance in the local administrator field.

Option B uses an outer opaque of type 8 and inter opaque of type 1 (see [Table 8: Recursive opaque types](#)).

[Figure 21: Non-segmented mLDP PMSI establishment \(Option B\)](#) depicts the processing required for I-PMSI and S-PMSI Inter-AS establishment.

Figure 21: Non-segmented mLDP PMSI establishment (Option B)



For non-segmented mLDP trees, A-D procedures follow those of the Intra-AS model, with the exception that NO EXPORT community must be excluded; LSP FEC includes mLDP VPN-recursive FEC.

For I-PMSI on Inter-AS Option B:

- A-D routes must be installed by ASBRs and next-hop information is changed as the routes are propagated, as shown in [Figure 21: Non-segmented mLDP PMSI establishment \(Option B\)](#).
- PMSI A-D routes are used to provide inter-domain connectivity on remote ASBRs.

On a receipt of an Intra-AS PMSI A-D route, PE2 resolves PE1's address (next-hop in PMSI route) to a labeled BGP route with a next-hop of ASBR3, because PE1 (Root-1) is not known via IGP. Because ASBR3 is not the originator of the PMSI route, PE2 sources an mLDP VPN recursive FEC with a root node of ASBR3, and an opaque value containing the information advertised by Root-1 (PE-1) in the PMSI A-D route, shown below, and forwards the FEC to ASBR 3 using IGP.

PE-2 LEAF FEC: (Root ASBR3, Opaque value {Root: ROOT-1, RD 60:60, Opaque Value: P2MPLSP-ID xx})

When the mLDP VPN-recursive FEC arrives at ASBR3, it notes that it is the identified root node, and that the opaque value is a VPN-recursive opaque value. Because Root-1 PE1 is not known via IGP, ASBR3 resolves the root node of the VPN-Recursive FEC using PMSI A-D (I or S) matching the information in the VPN-recursive FEC (the originator being PE1 (Root-1), RD being 60:60, and P2MP LSP ID xx). This yields ASBR1 as next hop. ASBR3 creates a new mLDP FEC element with a root node of ASBR1, and an opaque value being the received recursive opaque value, as shown below. ASBR then forwards the FEC using IGP.

ASBR-3 FEC: {Root ASBR 1, Opaque Value {Root: ROOT-1, RD 60:60, Opaque Value: P2MPLSP-ID xx}}

When the mLDP FEC arrives at ASBR1, it notes that it is the root node and that the opaque value is a VPN-recursive opaque value. As PE1's ROOT-1 address is known to ASBR1 through the IGP, no further recursion is required. Regular processing begins, using received Opaque mLDP FEC information.

ASBR-1 FEC: {Root: ROOT-1, Opaque Value: P2MP LSP-ID xx}



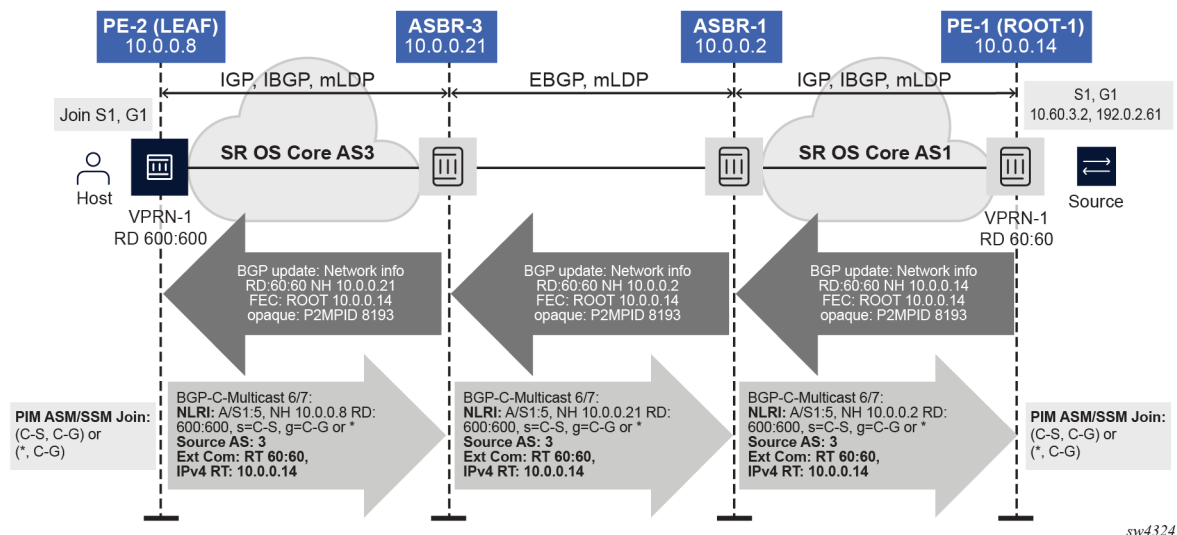
Note: VPN-Recursive FEC carries P2MPLSP ID. The P2MPLSP ID is used in addition to PE RD and Root to select a route to the mLDP root using the correct I-PMSI or S-PMSI route.

The functionality as described above for I-PMSI applies also to S-PMSI and (C-*, C-*) S-PMSI.

C-multicast Route Processing

C-multicast route processing functionality follows RFC 6513 section 8.1.2 (BGP used for route exchange). The processing is analogous to BGP Unicast VPN route exchange described in [Figure 49: Unicast VPN Option B with segmented MPLS](#) and [Figure 50: Unicast VPN Option C with segmented MPLS](#). [Figure 22: Non-segmented mLDP C-multicast exchange \(Option B\)](#) shows C-multicast route processing with non-segmented mLDP PMSI details.

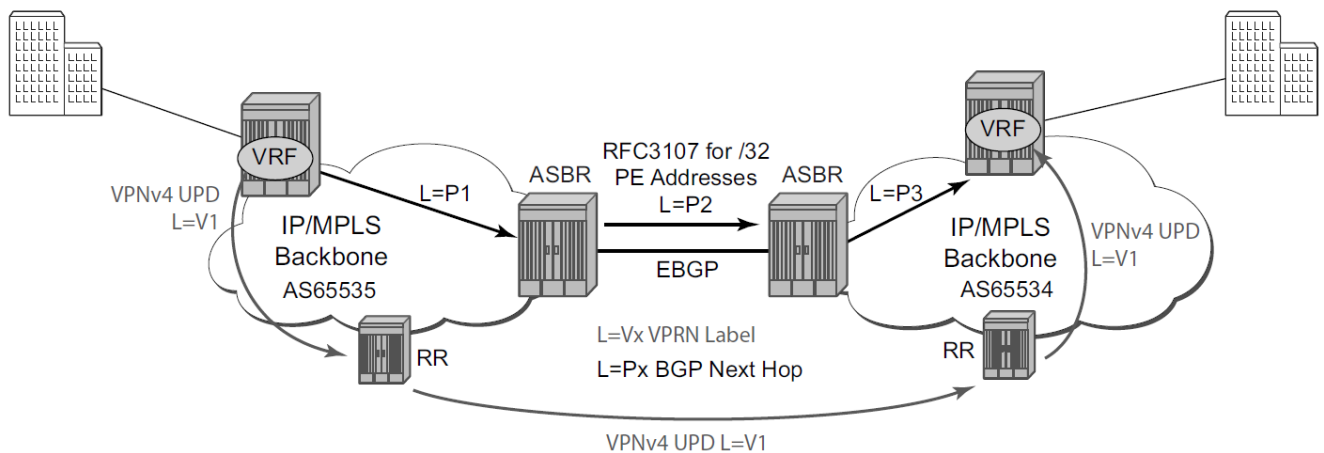
Figure 22: Non-segmented mLDP C-multicast exchange (Option B)



3.2.11.3 Inter-AS option-C

Inter-AS option-C allows a higher scale of VPRNs across AS boundaries and expands the trust model between the ASNs. As a result, this model is typically used within a single organization that may have multiple ASNs. The following figure shows inter-AS option-C.

Figure 23: Inter-AS option-C



OSSG257

The inter-AS option-C model differs from option-B, in that all direct knowledge of the remote AS is contained in and limited to the ASBR in an option-B network. The ASBR performs all necessary mapping functions, and the PE routers do not need to perform additional functions other than those performed by a non-inter-AS VPRN.

In option-C, however, knowledge from the remote AS is distributed throughout the local AS. This distribution allows for higher scalability, but also requires that all PEs and ASBRs involved in the inter-AS VPRNs participate in the exchange of inter-AS routing information. The ASBRs distribute reachability information for remote PE system IP addresses only by exchanging MP-eBGP labeled routes, as defined in RFC 8277, *Using BGP to Bind MPLS Labels to Address Prefixes*.

In option-C, the ASBRs distribute reachability information for remote PE's system IP addresses only. This is done between the ASBRs by exchanging MP-EBGP labeled routes, using RFC 8277, *Using BGP to Bind MPLS Labels to Address Prefixes*. Either RSVP-TE or LDP LSP can be selected to resolve next-hop for multihop EBGP peering by using the following command.

```
configure router bgp next-hop-resolution labeled-routes transport-tunnel
```

Distribution of VPRN routing information is handled by either direct MP-BGP peering between PEs in the different ASNs, or by one or more route reflectors in the ASN.

3.2.12 VPRN label security at inter-AS boundary

This feature allows the user to enforce security at an inter-AS boundary and to configure a router, acting in a PE role or in an ASBR role, or both, to accept packets of VPRN prefixes only from direct EBGP neighbors to which it advertised a VPRN label.

3.2.12.1 Feature configuration

To use this feature, network IP interfaces that can have the feature enabled must first be identified. Participating interfaces are identified as having the untrusted state. The router supports a maximum of 15 network interfaces that can participate in this feature.

Use the following command to configure the state of untrusted for a network IP interface.

```
configure router interface untrusted
```

Normally, the user applies the **untrusted** command to an inter-AS interface and PIP keeps track of the untrusted status of each interface. In the datapath, an inter-AS interface that is flagged by PIP causes the default forwarding to be set to the value of the **default-forwarding** option (**forward** or **drop**).

For backward compatibility, **default-forwarding** on the interface is set to the **forward** command option. This means that labeled packets are checked in the normal way against the table of programmed ILMs to decide if it should be dropped or forwarded in a GRT, a VRF, or a Layer 2 service context.

If the user sets the **default-forwarding** argument to the **drop** command option, all labeled packets received on that interface are dropped. For details, see [Datapath forwarding behavior](#).

This feature sets the default behavior for an untrusted interface in the data path and for all ILMs. To allow the data path to provide an exception to the normal way of forwarding handling away from the default for VPRN ILMs, BGP must flag those ILMs to the data path.

Use the following command to enable the exceptional ILM forwarding behavior, on a per-VPN-family basis.

```
configure router bgp neighbor-trust vpn-ipv4  
configure router bgp neighbor-trust vpn-ipv6
```

At a high level, BGP tracks each direct EBGP neighbor over an untrusted interface and to which it sent a VPRN prefix label. For each of those VPRN prefixes, BGP programs a bit map in the ILM that indicates, on a per-untrusted interface basis, whether the matching packets must be forwarded or dropped. For details, see [CPM behavior](#).

3.2.12.2 CPM behavior

This feature affects PIP behavior for management of network IP interfaces and in BGP for the resolution of BGP VPN-IPv4 and VPN-IPv6 prefixes.

The following are characteristics of CPM behavior related to PIP and the VPRN label security at inter-AS boundary feature:

- PIP manages the status of an untrusted interface based on the user configuration on the interface, as described in [Feature configuration](#). It programs the interface record in the data path using a 4-bit untrusted interface identification number. A trusted interface has no untrusted record.
- BGP determines the status of trusted or untrusted of an EBGP neighbor by checking the untrusted record provided by PIP for the index of the interface used by the EBGP session to the neighbor.
- BGP only tracks the status of trusted or untrusted for directly connected EBGP neighbors. The neighbor address and the local address must be on the same local subnet.
- BGP includes the neighbor status of trusted or untrusted in the tribe criteria. For example, if a group consists of two untrusted EBGP neighbors and one trusted EBGP neighbor and all three neighbors have the same local-AS, neighbor-AS, and export policy, then the result is two different tribes.

As a result, if the interface status changes from trusted to untrusted or untrusted to trusted, the EBGP neighbors on that interface bounce.

- When the feature is enabled for a specified VPN family and BGP advertises a label for one or more resolved VPN prefixes to a group of trusted and untrusted EBGP neighbors, it creates a 16-bit map

in the ILM record in which it sets the bit position corresponding to the identification number of each untrusted interface used by a EBGp session to a neighbor to which it sent the label.

A bit in the ILM record bit-map is referred to as the untrusted interface forwarding bit. The bit position corresponding to the identification number of any other untrusted interface is left clear.

For details on the data path of the ILM bit-map record, see [Datapath forwarding behavior](#).

- Because the same label value is advertised for prefixes in the same VRF (label per-VRF mode) and for prefixes with the same next hop (label per-next-hop mode), BGP programs the forwarding bit position in the ILM bit map for both VPN IPv4 and VPN IPv6 prefixes sharing the same label, as long as the feature is enabled for at least one of the two VPN families.
- BGP tracks, on a per-untrusted interface basis, the number of RIB-Out entries to EBGp neighbors that reference a specific VPN label. When that reference transitions from zero to a positive value or from a positive value to zero, the label for the ILM of the VPN prefix is re-downloaded to the IOM with the forwarding bit position in the ILM bit map record updated accordingly (set or unset, respectively).

This feature supports label per-VRF and label per-next-hop modes for the PE role. The feature supports label per-next-hop mode for the ASBR role.

The feature is not supported with label per-prefix mode in a PE role and is not supported in a Carrier Supporting Carrier (CSC) PE role.

3.2.12.3 Datapath forwarding behavior

ILM forwarding on a trusted interface behaves as in earlier releases and is not changed. The ILM forwarding bit map is ignored and packets are forwarded normally.

ILM forwarding on an untrusted interface follows these rules:

- Only the top-most label in the label stack in a received packet is checked against the next set of rules. The top label can correspond to any one of the following applications:
 - a transport label with a pop or swap operation of static, RSVP-TE, SR-TE, LDP, SR-ISIS, SR-OSPF, or BGP-LU
 - a BGP VPRN inter-AS option B label with a swap operation when the router acts in the ASBR role for VPN routes
 - a service delimiting label for a local VRF when the router acts as a PE in a VPRN service
- The datapath checks the bit position in the bit map in the ILM record, when present, that corresponds to the untrusted interface identification number in the interface record and then makes a forwarding decision to drop or forward.

A decision to forward means that a labeled packet proceeds to the regular ILM processing and its label stack is checked against the table of programmed ILMs to decide if the packet should be:

- dropped
- forwarded to CPM
- forwarded as an MPLS packet
- forwarded as an IP packet in a GRT or a VRF context
- forwarded as a packet in a Layer 2 service context
- The following are the processing rules of the ILM:
 - interface **default-forwarding** = **forward** and ILM bit-map not present ⇒ forward packet

- interface **default-forwarding** = **forward** and interface forwarding bit position in the ILM bit-map 1 ⇒ forward packet
 - interface **default-forwarding** = **forward** and interface forwarding bit position in the ILM bit-map zero ⇒ drop packet
 - interface **default-forwarding** = **drop** and ILM bit-map not present ⇒ drop packet
 - interface **default-forwarding** = **drop** and interface forwarding bit position in the ILM bit-map zero ⇒ drop packet
 - interface **default-forwarding** = **drop** and interface forwarding bit position in the ILM bit-map 1 ⇒ forward packet
 - When the EBGP neighbor is not directly connected, BGP does not track that neighbor (see [CPM behavior](#)). In this case, the VPRN packet is received with a transport label or without a transport label if implicit-null is enabled in LDP or RSVP-TE for the transport label. Either way, the forwarding decision for the packet is solely dictated by the configuration of the **default-forwarding** command option on the incoming interface.
 - If the direct EBGP neighbor sends a VPRN packet using the MPLS-over-GRE encapsulation, the datapath does not check the interface forwarding bit position in the ILM bit map. In this case, the forwarding decision of the packet is solely dictated by the configuration of the **default-forwarding** command option on the incoming interface.
- SR OS EBGP neighbors never use the MPLS-over-GRE encapsulation over an inter-AS link, but third party implementations may do this.

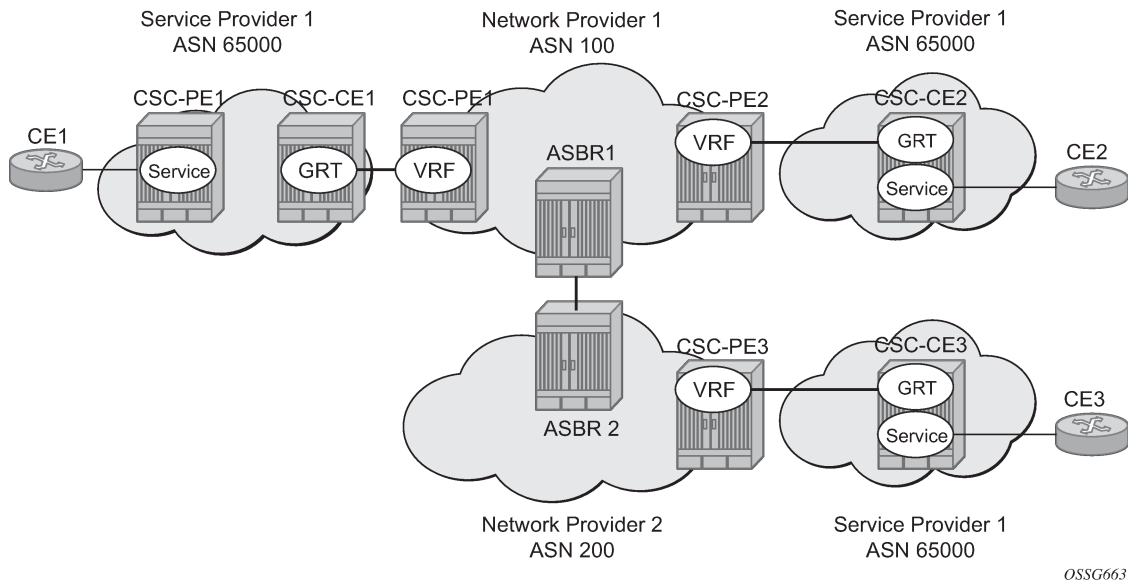
3.2.13 CSC

Carrier Supporting Carrier (CSC) is a solution that allows one service provider (the Customer Carrier) to use the IP VPN service of another service provider (the Super Carrier) for some or all of its backbone transport. RFC 4364 defines a Carrier Supporting Carrier solution for BGP/MPLS IP VPNs that uses MPLS on the interconnections between the two service providers to provide a scalable and secure solution.

CSC support in SR OS allows the 7705 SAR Gen 2 to be deployed as any of the following devices shown in [Figure 24: Carrier Supporting Carrier reference diagram](#):

- PE1 (service provider PE)
- CSC-CE1, CSC-CE2 and CSC-CE3 (CE device from the point of view of the backbone service provider)
- CSC-PE1, CSC-PE2 and CSC-PE3 (PE device of the backbone service provider)
- ASBR1 and ASBR2 (ASBR of the backbone service provider)

Figure 24: Carrier Supporting Carrier reference diagram



3.2.13.1 Terminology

CE	customer premises equipment dedicated to one particular business/enterprise
PE	service provider router that connects to a CE to provide a business VPN service to the associated business/enterprise
CSC-CE	an ASBR/peering router that is connected to the CSC-PE of another service provider for purposes of using the associated CsC IP VPN service for backbone transport
CSC-PE	a PE router belonging to the backbone service provider that supports one or more CSC IP VPN services

3.2.13.2 CSC connectivity models

A PE router deployed by a customer service provider to provide Internet access, IP VPNs, or L2 VPNs may connect directly to a CSC-PE device, or it may back haul traffic within its local "site" to the CSC-CE that provides this direct connection. Here, "site" means a set of routers owned and managed by the customer service provider that can exchange traffic through means other than the CSC service. The function of the CSC service is to provide IP/MPLS reachability between isolated sites.

The CSC-CE is a "CE" from the perspective of the backbone service provider. There may be multiple CSC-CEs at a specific customer service provider site and each one may connect to multiple CSC-PE devices for resiliency/multihoming purposes.

The CSC-PE is owned and managed by the backbone service provider and provides CSC IP VPN service to connected CSC-CE devices. In many cases, the CSC-PE also supports other services, including regular business IP VPN services. A single CSC-PE may support multiple CSC IP VPN services. Each customer

service provider is allocated its own VRF within the CSC-PE; VRFs maintain routing and forwarding separation and allow the use of overlapping IP addresses by different customer service providers.

A backbone service provider may not have the geographic span to connect, with reasonable cost, to every site of a customer service provider. In this case, multiple backbone service providers may coordinate to provide an inter-AS CSC service. Different inter-AS connectivity options are possible, depending on the trust relationships between the different backbone service providers.

3.2.13.3 CSC-PE configuration and operation

This section applies to CSC-PE1, CSC-PE2 and CSC-PE3 in [Figure 24: Carrier Supporting Carrier reference diagram](#).

3.2.13.4 CSC interface

From the point of view of the CSC-PE, the IP/MPLS interface between the CSC-PE and a CSC-CE has these characteristics:

1. The CSC interface is associated with one (and only one) VPRN service. Routes with the CSC interface as next-hop are installed only in the routing table of the associated VPRN.
2. The CSC interface supports EBGP or IBGP for exchanging labeled IPv4 routes (RFC 8277). The BGP session may be established between the interface addresses of the two routers or else between a loopback address of the CSC-PE VRF and a loopback address of the CSC-CE. In the latter case, the BGP next-hop is resolved by either a static or OSPFv2 route.
3. An MPLS packet received on a CSC interface is dropped if the top-most label was not advertised over a BGP (RFC 8277) session associated with one of the VPRN's CSC interfaces.
4. The CSC interface supports ingress QoS classification based on 802.1p or MPLS EXP. It is possible to configure a default FC and default profile for the CSC interface.
5. The CSC interface supports QoS (re)marking for egress traffic. Policies to remark 802.1p or MPLS EXP based on forwarding-class and profile are configurable per CSC interface.
6. By associating a port-based egress queue group instance with a CSC interface, the egress traffic can be scheduled/shaped with per-interface, per-forwarding-class granularity.
7. By associating a forwarding-plane based ingress queue group instance with a CSC interface, the ingress traffic can be policed to per-interface, per-forwarding-class granularity.
8. Ingress and egress statistics and accounting are available per CSC interface. The exact set of collected statistics depends on whether a queue-group is associated with the CSC interface, the traffic direction (ingress vs. egress), and the stats mode of the queue-group policers.

An Ethernet port or LAG with a CSC interface can be configured in hybrid mode or network mode. The port or LAG supports null, dot1q, or QinQ encapsulation. Use the following commands to create a CSC interface on a port or LAG in null mode.

```
configure service vprn network-interface port port-id
configure service vprn network-interface lag lag-id
```

Use the following commands to create a CSC interface on a port or LAG in dot1q mode.

```
configure service vprn network-interface port port-id:qtag1
configure service vprn network-interface lag lag-id:qtag1
```

Use the following commands to create a CSC interface on a port or LAG in QinQ mode.

```
configure service vprn network-interface port port-id:qtag1.qtag2
configure service vprn network-interface port port-id:qtag1.*
configure service vprn network-interface lag lag-id:qtag1.qtag2
configure service vprn network-interface lag lag-id:qtag1.*
```

A CSC interface supports the same capabilities (and supports the same commands) as a Base router network interface, except it does not support:

- IPv6
- LDP
- RSVP
- Proxy ARP (local/remote)
- Network domain configuration
- DHCP
- Ethernet CFM
- Unnumbered interfaces

3.2.13.5 QoS

3.2.13.5.1 Egress

Egress traffic on a CSC interface can be shaped and scheduled by associating a port-based egress queue-group instance with the CSC interface. The steps for doing this are summarized below:

Procedure

- Step 1.** Create an egress queue-group-template.
- Step 2.** Define one or more queues in the egress queue-group. For each one specify scheduling command options such as CIR, PIR, CBS and MBS and, if using H-QoS, the parent scheduler.
- Step 3.** Apply an instance of the egress queue-group template to the network egress context of the Ethernet port with the CSC interface. When doing so, and if applicable, associate an accounting policy or a scheduler policy, or both, with this instance.
- Step 4.** Create a network QoS policy.
- Step 5.** In the egress part of the network QoS policy define EXP remarking rules, if necessary.
- Step 6.** In the egress part of the network QoS policy map a forwarding class to a queue ID using the **port-redirect-group** command.

Example

MD-CLI

```
[ex:/configure qos network "2" egress fc l2]
A:admin@node-2# info
    port-redirect-group {
        queue 5
    }
```

Example

Classic CLI

```
A:node-2>config>qos>network>egress$ info
-----
fc l2
  port-redirect-group queue 5
exit
-----
```

- Step 7.** Apply the network QoS policy created in step 4 to the CSC interface and specify the name of the egress queue-group created in step 1 and the specific instance defined in step 3.

3.2.13.5.2 Ingress

Ingress traffic on a CSC interface can be policed by associating a forwarding-plane based ingress queue-group instance with the CSC interface. The steps for doing this are summarized below:

Procedure

- Step 1.** Create an ingress queue-group-template.
- Step 2.** Define one or more policers in the ingress queue-group. For each one specify command options such as CIR, PIR, CBS and MBS and, if using H-Pol, the parent arbiter.
- Step 3.** Apply an instance of the ingress queue-group template to the network ingress context of the forwarding plane with the CSC interface.
When doing so, and if applicable, associate an accounting policy or a policer-control-policy, or both, with this instance.
- Step 4.** Create a network QoS policy.
- Step 5.** In the ingress part of the network QoS policy define EXP classification rules, if necessary.
- Step 6.** In the ingress part of the network QoS policy map a forwarding class to a policer ID using the **fp-redirect-group policer** command.

Example

MC-CLI

```
[ex:/configure qos network "3" ingress fc l2]
A:admin@node-2# info
  fp-redirect-group {
    policer 5
  }
```

Example

classic CLI

```
A:node-2>config>qos>network>ingress$ info
-----
fc l2
  fp-redirect-group policer 5
exit
-----
```

- Step 7.** Apply the network QoS policy created in step 4 to the CSC interface and specify the name of the ingress queue-group created in step 1 and the specific instance defined in step 3.

3.2.13.6 MPLS

BGP-8277 is used as the label distribution protocol on the CSC interface. When BGP in a CSC VPRN needs to distribute a label corresponding to a received VPN-IPv4 route, it takes the label from the global label space. The allocated label is not re-used for any other FEC regardless of the routing instance (Base router or VPRN). If a label L is advertised to the BGP peers of CSC VPRN A then a received packet with label L as the top most label is only valid if received on an interface of VPRN A, otherwise the packet is discarded.

To use BGP-8277 as the label distribution protocol on the CSC interface, add the **family label-ipv4** command to the family configuration at the instance, group, or neighbor level. This causes the capability to send and receive labeled-IPv4 routes {AFI=1, SAFI=4} to be negotiated with the CSC-CE peers.

3.2.13.7 CSC VPRN service configuration

To configure a VPRN to support CSC service, the **carrier-carrier-vpn** command must be enabled. The command fails if the VPRN service has any existing SAP or spoke SDP interfaces. A CSC interface can be added to a VPRN (using the **network-interface** command) only if the **carrier-carrier-vpn** command is enabled.

A VPRN service with the **carrier-carrier-vpn** command may be provisioned to use **auto-bind-tunnel**, configured spoke SDPs, or some combination. All SDP types are supported except for:

- GRE SDPs
- LDP over RSVP-TE tunnel SDPs

Other aspects of VPRN configuration are the same in a CSC model as in a non-CSC model.

3.2.14 Node management using VPRN

There are two basic approaches that can be used to manage a node using a VPRN. In both cases, management traffic is received and sent in the VPRN router instance:

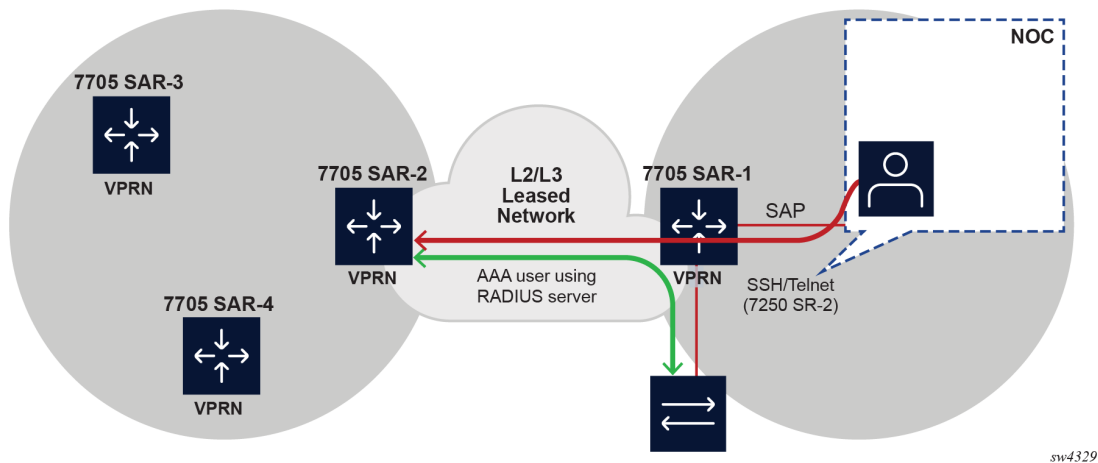
- Management traffic can target the IP address of a VPRN interface
- Management traffic can target the system address in the Base router instance (using GRT leaking)

In the first approach, node management can be enabled using the local interface of any VPRN service. A management VPRN is separated from other traffic using an MPLS transport tunnel. This provides IP domain separation and security for the management domain. The management domain supports IPv4 and IPv6 address families and the AAA server is connected to the same VPRN for authentication, authorization, and accounting. The SR OS allows management using a VPRN as long the management packet is destined for a local interface of the VPRN, in addition it allows configuration of the AAA servers within a VPRN; see [Figure 25: VRF network example](#).

In the second approach, node management is achieved using GRT leaking. In this case the management traffic uses an IP address in the Base routing context. See [Management via VPRN using GRT leaking](#) for details on this method.

The remainder of this section describes node management using the local VPRN interfaces (non GRT leaking).

Figure 25: VRF network example



3.2.14.1 VPRN management

VPRN management can be enabled by configuring the appropriate management protocol within the VPRN from the following context.

```
configure service vprn management
```

The following protocols can be enabled in this context:

- FTP
- gRPC
- NETCONF
- SSH
- Telnet
- Telnetv6

SNMP access is not controlled in the following context. See section [SNMP management](#).

```
configure service vprn management
```

NTP and PTP access are not controlled in the following context. See [NTP within a VPRN service](#) and [PTP within a VPRN service](#).

```
configure service vprn management
```

By default, all protocols are disabled. When one of these protocols is enabled, that VRF becomes a management VRF.



Note: An SSH server can only be enabled or disabled in VPRN. All other SSH configurations such as key re-exchange, SSH client/server HMAC/cipher list, and so on, are configured under the system and used globally.

All other gRPC configurations remain global under the following context.

```
configure system grpc
```

All other NETCONF configurations remain global under the following context.

```
configure system management-interface netconf
```

Configure the TLS profiles needed for gRPC under the following context.

```
configure system security tls
```

3.2.14.2 AAA management

Use the commands in the following context to configure the authentication, authorization, and accounting order for the system, including VPRNs:

- **MD-CLI**

```
configure system security user-params authentication-order order
```

- **classic CLI**

```
configure system security password authentication-order
```

Use the commands in the following context to configure the system local user profile configuration for local user authentication and authorization, including VPRNs:

- **MD-CLI**

```
configure system security aaa local-profiles profile
```

- **classic CLI**

```
configure system security profile
```

Use the commands in the following contexts to configure AAA servers:

- **MD-CLI**

- **System AAA servers**

```
configure system security aaa
```

- **AAA remote servers under the VPRN**

```
configure service vprn aaa remote-servers
```

- **classic CLI**

- **System AAA servers**

```
configure system security
```

- **AAA remote servers under the VPRN**

```
configure service vprn aaa remote-servers
```

When AAA servers are configured using the preceding commands, they are used as follows:

- If servers are configured under the VPRN AAA, only the VPRN AAA servers are used.
For example, the **authentication-order** command lists the order as local, TACACS+, and RADIUS, while the VPRN only has a RADIUS server configured, and under the system AAA servers both TACACS+ and RADIUS are configured. In this case, if a management session connects to the VPRN and the destination IP matches a local interface in the VPRN, the SR OS tries the local AAA first, and then RADIUS as configured in the VPRN. The SR OS does not try the system AAA servers because there is a AAA server configured in the VPRN.
- If servers are configured under VPRN AAA and the VPRN AAA command options are configured for in-band, out-of-band, or VPRN, the servers can be used for the VPRN and the system.
- If no AAA servers are configured under VPRN AAA, the system AAA servers are used.

3.2.14.3 SNMP management

The SR OS SNMP agent can be reached via a VPRN interface address when the following command is enabled:

- **MD-CLI**

```
configure service vprn snmp access true
```

- **classic CLI**

```
configure service vprn snmp access
```

Using an SNMP community defined inside the VPRN context or a user associated with an SNMPv3 USM access group defined in the system context allows access to a subset of the full SNMP data model.

Use the following command to define an SNMP community inside the VPRN context.

```
configure service vprn snmp community
```

Use the following command to define a user associated with an SNMPv3 USM access group in the system context.

```
configure system snmp access
```

Use the following command to view this subset.

```
show system security view
```

Use an SNMP community defined in the system context to allow access to the full SNMP data model (unless otherwise restricted used SNMP views). Use the following command to create the SNMP community strings for SNMPv1 and SNMPv2 access.

```
configure system security snmp community
```

Alternatively, GRT leaking and a Base routing IP address can be used (along with an SNMP community defined at the **system** context) to allow access to the entire SNMP data model (see the **allow-local-management** command).

A network manager using SNMP, cannot discover or fully manage an SR OS router using an SNMP community defined inside the VPRN context. Full SNMP access requires using one of the approaches described above.

SNMP communities configured under a VPRN are associated with the SNMP context "vprn". For example, walking the ifTable (IF-MIB) using the community configured for VPRN 5 returns counters and status for interfaces in VPRN 5 only.



Note: To access the Base router ifTable entries in a VPRN, use the community string that is defined in the system context.

```
configure system security snmp community
```

To access VPRN ifTable entries, use the community string that is defined inside that VPRN context.

```
configure service vprn snmp community
```

3.2.14.4 Events and notifications

Syslog, SNMP traps, and NETCONF notifications are generated via the Event Logging System.

Use the commands in the following context to define the VPRN syslog destinations.

```
configure service vprn log syslog
```

Use the commands in the following context to define the SNMP trap destination.

```
configure service vprn log snmp-trap-group
```

Use the following command to direct events for the whole system to a destination within the management VPRN.

```
configure log services-all-events
```

See the *7705 SAR Gen 2 System Management Guide* for more information about this command.

3.2.14.5 DNS resolution

DNS default domain name and DNS servers for domain name resolution can be defined within a VPRN in the following context.

```
configure service vprn dns
```

3.2.15 Traffic leaking to GRT

Traffic leaking to Global Route Table (GRT) for the 7705 SAR Gen 2 allows service providers to offer VPRN and Internet services to their customers over a single VRF interface.

Packets entering a local VRF interface can have route processing results derived from the VPRN forwarding table or the GRT. The leaking and preferred lookup results are configured on a per VPRN basis. Configuration can be general (for example, any lookup miss in the VPRN forwarding table can be resolved in the GRT), or specific (for example, specific routes should only be looked up in the GRT and ignored in the VPRN). To provide operational simplicity and improve streamlining, the CLI configuration is contained within the context of the VPRN service.

Use the commands in following context to configure the traffic leaking to GRT feature:

- **MD-CLI**

In the MD-CLI, configuring **grt-lookup** to **true** enables the basic functionality.

```
configure service vprn grt-leaking
```

- **classic CLI**

In the classic CLI, the **enable-grt** command establishes the basic functionality.

```
configure service vprn grt-lookup
```

This is an administrative context and provides the container under which the user can enter all specific commands, except policy definition. Policy definitions remain unchanged but are referenced from this context.

When it is configured, any lookup miss in the VRF table is resolved in the GRT, if available. By itself, this only provides part of the solution. Packet forwarding within GRT must route packets back to the correct node and to the specific VPRN from which the destination exists. Destination prefixes must be leaked from the VPRN to the GRT through the use of policy. Use the commands in the following context to create the policies:

- **MD-CLI**

```
configure policy-options
```

- **classic CLI**

```
configure router policy-options
```

By default, the number of prefixes leaked from the VPRN to the GRT is limited to five. Use the following command to override the default or remove the limit:

- **MD-CLI**

```
configure service vprn grt-leaking export-limit
```

- **classic CLI**

```
configure service vprn grt-lookup export-limit
```

When a VPRN forwarding table consists of a default route or an aggregate route, the customer may require the service provider to poke holes in those, or provide more specific route resolution in the GRT. In this case, the service provider may configure a static-route-entry and specify the GRT as the nexthop type.

The lookup result prefers any successful lookup in the GRT that is equal to or more specific than the static route, bypassing any successful lookup in the local VPRN.

This feature and Unicast Reverse Path Forwarding (uRPF) are mutually exclusive. When a VPRN service is configured with either of these functions, the other cannot be enabled. The following type of routes are not leaked from VPRN into the Global Routing Table (GRT):

- Aggregate routes
- BGP VPN extranet routes

3.2.15.1 Management via VPRN using GRT leaking

In addition to node management using the IP addresses of VPRN interfaces, see [Node management using VPRN](#), management via a VPRN can also be achieved using IP addresses in the Base routing instance and GRT leaking.

When a management packet arrives on a VPRN, there is a lookup for the destination IP address of the packet. If the destination IP is resolved using VPRN and the corresponding protocol is enabled under VPRN management, then the packet is extracted to CPM.

If the destination IP address is not a VRF IP and GRT leaking is enabled, a second lookup is done in the GRT FIB. If the IP address belongs to a local interface in GRT and **allow-local-management** is enabled under the following context, the packet is extracted using GRT leaking to the CPM.

- **MD-CLI**

```
configure service vprn grt-leaking
```

- **classic CLI**

```
configure service vprn grt-lookup enable-grt
```

3.2.16 Traffic leaking from VPRN to GRT for IPv6

This feature allows IPv6 destination lookups in two distinct routing tables. IPv6 packets within a Virtual Private Routed Network (VPRN) service is able to perform a lookup for IPv6 destination against the Global Route Table (GRT) as well as within the local VPRN.

Currently, VPRN to VPRN routing exchange is accomplished through the use of import and export policies based on Route Targets (RTs), the creation of extranets. This new feature allows the use of a single VPRN interface for both corporate VPRN routing and other services (for example, Internet) that are reachable

outside the local routing context. This feature takes advantage of the dual lookup capabilities in two separate routing tables in parallel.

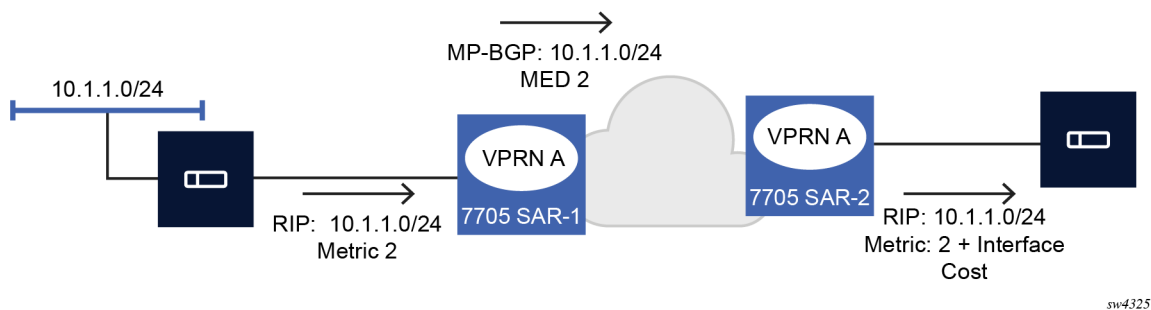
This feature enables IPv6 capability in addition to the existing IPv4 VPRN-to-GRT Route Leaking feature.

3.2.17 RIP metric propagation in VPRNs

When RIP is used as the PE-CE protocol for VPRNs (IP-VPNs), the RIP metric is only used by the local node running RIP with the Customer Equipment (CE). The metric is not used to or encoded into and MP-BGP path attributes exchanged between PE routers.

The RIP metric can also be used to be exchanged between PE routers so if a customer network is dual homed to separate PEs the RIP metric learned from the CE router can be used to choose the best route to the destination subnet. By using the learned RIP metric to set the BGP MED attribute, remote PEs can choose the lowest MED and in turn the PE with the lowest advertised RIP metric as the preferred egress point for the VPRN. The following figure shows RIP metric propagation in VPRNs.

Figure 26: RIP metric propagation in VPRNs



3.2.18 NTP within a VPRN service

Communication to external NTP clocks through VPRNs is supported in two ways: communication with external servers and peers, and communication with external clients.

Communication with external servers and peers are controlled using the same commands as used for access via base routing (see Network Time Protocol (NTP) in the *7705 SAR Gen 2 Basic System Configuration Guide*).

Communication with external clients is controlled via the VPRN routing configuration. The support for the external clients can be as unicast or broadcast service. In addition, authentication keys for external clients are configurable on a per-VPRN basis.

Only a single instance of NTP remains in the node that is time sourced to as many as five NTP servers attached to the base or management network.

The NTP **show** command displays NTP servers and all known clients. Because NTP is UDP-based only, no state is maintained. As a result, the **show** command output only displays when the last message from the client was received.

3.2.19 PTP within a VPRN service

The PTP within a VPRN service provides access to the PTP clock within the 7705 SAR Gen 2 through one or more VPRN services. Only one VPRN or the base routing instance may have configured peers, but all may have discovered peers. If needed, a limit on the maximum number of dynamic peers allowed may be configured on a per routing instance basis.

For more information about PTP see the *7705 SAR Gen 2 Basic System Configuration Guide*.

3.2.20 VPN route label allocation

The method used for allocating a label value to an originated VPN-IP route (exported from a VPRN) depends on the configuration of the VPRN service and its VRF export policies. SR OS supports three label modes:

- label per VRF
- label per next hop (LPN)
- label per prefix (LPP)

Label per VRF is the label allocation default. It is used when the label mode is configured as VRF (or not configured) and the VRF export policies do not apply an **advertise-label per-prefix** action. All routes exported from the VPRN with the per-VRF label have the same label value. When the PE receives a terminating MPLS packet with a per-VRF label, the label value selects the VRF context in which to perform a forwarding table lookup and this lookup determines the outgoing interface (or set of interfaces if ECMP applies).

Label per next hop is used when the exported route is not a local or aggregate route, the label mode is configured as next-hop, and the VRF export policies do not apply an advertise-label per-prefix override. It is also used when an inactive (backup path) BGP route is exported by effect of the **export-inactive-bgp** command if there is no advertise-label per-prefix override. All LPN-exported routes with the same primary next hop have the same per-next-hop label value. When the PE receives a terminating MPLS packet with a per-next-hop label, the label lookup selects the outgoing interface for forwarding, without any FIB lookup that may cause problems with overlapping prefixes. LPN does not support ECMP, BGP fast reroute, QPPB, or policy accounting features that may otherwise apply.



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

Label per-prefix is used when a qualifying IP route is exported by matching a VRF export policy action with **advertise-label per-prefix**. Any IPv4 or IPv6 route that is not a local route, aggregate route, BGP-VPN route, or GRT lookup static route qualifies. With LPP, every prefix is associated with its own unique label value that does not change while the route is present in the route table. When the PE receives a terminating MPLS packet with a per-prefix label value, the packet is forwarded as if the FIB lookup found only the matching prefix route and not any of the more specific prefix routes that would normally be selected. LPP supports ECMP, QPPB, and policy accounting as part of the egress forwarding decision. It does not support BGP fast reroute or BGP sticky ECMP.

The following points summarize the logic that determines the label allocation method for an exported route:

- If the IP route is LOCAL, AGGREGATE, or BGP-VPN always use the VRF label.
- If the IP route is accepted by a VRF export policy with the **advertise-label per-prefix** action, use LPP.

- If the IP (BGP) route is exported by the **export-inactive-bgp** command (VPRN best external), use LPN.
- If the IP route is exported by a VPRN configured for label-mode next-hop, use LPN.
- Else, use the per-VRF label.

3.2.20.1 Configuring the service label mode

Use the following command to change the service label mode of the VPRN for the 7705 SAR Gen 2.

```
configure service vprn label-mode
```

The default mode (if the command is not present in the VPRN configuration) is **vrf**, meaning distribution of one service label for all routes of the VPRN. When a VPRN X is configured with the **label-mode next-hop** command option, the service label that it distributes with an IPv4 or IPv6 route that it exports depends on the type of route as summarized in [Table 5: Service labels distributed in service label per next hop mode](#).

Table 5: Service labels distributed in service label per next hop mode

Route type	Distributed service label
remote route with IP A (associated with a SAP) as resolved next-hop	platform-wide unique label allocated to next-hop A
remote route with IP B (associated with a spoke SDP) as resolved next-hop	platform-wide unique label allocated to next-hop B
local route	platform-wide unique label allocated to VPRN X
aggregate route	platform-wide unique label allocated to VPRN X
ECMP route	platform-wide unique label allocated to next-hop A (the lowest next-hop address in the ECMP set)
BGP route with a backup next-hop (BGP FRR)	platform-wide unique label allocated to next-hop A (the lowest next-hop address of the primary next-hops)

In the classic CLI, a change to the label mode of a VPRN requires the VPRN to first be administratively disabled.

3.2.20.2 Restrictions and usage notes

The service label per next-hop mode has the following restrictions:

- **ECMP**

The VPRN label mode should be set to VRF if distribution of traffic across the multiple PE-CE next-hop interfaces of an ECMP route is needed.

- **hub and spoke VPN**

The VPRN label mode should not be set to next-hop if the user does not want the hub-connected CE to be involved in the forwarding of spoke-to-spoke traffic.

- **BGP next-hop indirection**

BGP next-hop indirection has no benefit in service label per next-hop mode. When the resolved next-hop interface of a BGP next-hop changes all of the affected BGP routes must be re-advertised to VPRN peers with the new service label corresponding to the new resolved next-hop.

- **BGP anycast**

When a PE failure results in redirection of MPLS packets to the other PE in a dual-homed pair, the service label mode is forced to VRF, for example, FIB lookup determines the next-hop even if the label mode of the VPRN is configured as next-hop.

- **U-turn routing**

U-turn routing is effectively disabled by service-label per next-hop.

- **Carrier Supporting Carrier**

The label-mode configuration of a VPRN with CSC interfaces is ignored for BGP-8277 routes learned from connected CSC-CE devices.

3.2.21 VPRN Support for BGP FlowSpec

When a VPRN BGP instance receives an IPv4 or IPv6 flow route, and that route is valid/best, the system attempts to construct an IPv4 or IPv6 filter entry from the NLRI contents and the actions encoded in the UPDATE message. If the attempt is successful, the filter entry is added to the system-created "fSpec-*n*" IPv4 or IPv6 embedded filter, where *n* is the service-id of the VPRN. These embedded filters may be inserted into configured IPv4 and IPv6 filter policies that are applied to ingress traffic on a selected set of the VPRN's IP interfaces. These interfaces can include SAP and spoke SDP interfaces, but not CsC network interfaces.

When FlowSpec rules are embedded into a user-defined filter policy, the insertion point of the rules is configurable through the **offset** command option in the following contexts:

- **MD-CLI**

```
configure filter ip-filter embed filter
configure filter ipv6-filter embed filter
```

- **classic CLI**

```
configure filter ip-filter embed-filter
configure filter ipv6-filter embed-filter
```

The sum of the **ip-filter-max-size** and **offset** must not exceed the maximum filter entry ID range.

3.2.22 MPLS EL and hash label

The router supports both the MPLS EL (RFC 6790) and the FAT label, known as the hash label (RFC 6391). 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 EL is supported for VPRN services, as well as Epipe spoke-SDP termination on VPRN interfaces. To configure insertion of the EL, use the **entropy-label** command in the **vprn** context or **spoke-sdp** context of an interface.

The hash label is also supported for Epipe spoke-SDP termination on VPRN and VPRN services bound to any MPLS-type encapsulated SDP, as well as to a VPRN service using the **auto-bind-tunnel** command with the **resolution-filter** configured as any MPLS tunnel type. Configure the hash label using the **hash-label** command in the following contexts.

```
configure service vprn
configure service vprn spoke-sdp
configure service vprn interface spoke-sdp
```

Either the hash label or the EL can be configured on one object, but not both.

3.2.23 LSP tagging for BGP next hops or prefixes and BGP-LU

It is possible to constrain the tunnels used by the system for resolution of BGP next-hops or prefixes and BGP labeled unicast routes using LSP administrative tags. See the "LSP Tagging and Auto-Bind Using Tag Information" section of the *7705 SAR Gen 2 MPLS Guide* for more information.

3.2.24 Route leaking from GRT to VPRN instances

The GRT to VPRN route leaking feature allows actual routes from the global route table to be exported into specific VPRN instances allowing those routes to be used for forwarding as well as being re-advertised within the VPRN context.

There are two stages for route leaking. The first stage requires the configuration of a set of leak-export route policies that identify which GRT routes are subject to being exported into VPRN services. The **leak-export** command in the **configure router** context is used to configure between one and four route policies. The GRT routes must match the policy entries configured with the following command to **accept**:

- **MD-CLI**

```
configure policy-options policy-statement entry action action-type
```

- **classic CLI**

```
configure router policy-options policy-statement entry action
```

In addition, the **leak-export-limit** command is used to specify the maximum number of GRT routes that can be included in the GRT leak pool.

The second stage requires the configuration of an import GRT policy that specifies which routes within the GRT leak pool that are leaked into the associated VPRN instances route table. The **import-grt** command in the following context is used and accepts one route policy:

- **MD-CLI**

```
configure service vprn grt-leaking
```

- **classic CLI**

```
configure service vprn grt-lookup
```

For the GRT route to be leaked into the local VPRN, the route must match a policy entry with following command set to **accept**:

- **MD-CLI**

```
configure policy-options policy-statement entry action action-type
```

- **classic CLI**

```
configure router policy-options policy-statement entry action
```

If a GRT route passes both stages, it is added into the VPRN route table which it to be used for IP forwarding as well as re-advertisement within other routing protocols within the VPRN context.

Both IPv4 and IPv6 routes can be leaked using this process from the GRT into one or more VPRN instances. The GRT route types that can be leaked using this process are:

- RIP, OSPF, and IS-IS routes
- Direct routes
- Static routes

3.2.25 Class-based forwarding of VPN-v4/v6 prefixes over RSVP-TE or SR-TE LSPs

This feature enables class-based forwarding (CBF) with ECMP of BGP VPN-v4/v6 prefixes that are resolved using RSVP-TE or SR-TE configured as **auto-bind-tunnel**.

3.2.25.1 Feature configuration

To configure this feature:

- Enable resolution to RSVP-TE or SR-TE tunnels in the **auto-bind-tunnel** context.
- Enable ECMP in the **auto-bind-tunnel** context.
- Enable class-forwarding in the **vprn** context.
- Define at least one class forwarding policy in the **mpls** context, the FC to sets associations and the LSP to (policy, set) associations.

The SR OS CBF implementation supports spraying of packets over a maximum of six forwarding sets of ECMP LSPs only when the system profile is **profile-b** that is supported on an FP4 or later-based CPM. In any other case, the maximum number of forwarding sets of ECMP LSPs is four.

Example: MD-CLI

```
[ex:/configure router "Base" mpls]
A:admin@node-2# info
  class-forwarding-policy "test" {
    fc l2 {
      forwarding-set 2
    }
    fc af {
```

```

        forwarding-set 3
    }
    fc l1 {
        forwarding-set 4
    }
}

```

Example: classic CLI

```

A:node-2>config>router>mpls$ info
-----
...
    class-forwarding-policy "test"
        fc l2 forwarding-set 2
        fc af forwarding-set 3
        fc l1 forwarding-set 4
    exit
-----

```

All FCs are mapped to set 1 as soon as the policy is created. The user can make changes to the mapping of FCs as required. An FC that is not added to the class-forwarding policy, is always mapped to set 1. At most, an FC can be mapped to a single forwarding set. One or more FCs can be mapped to the same set. The user can indicate the initial default set by including the *default-set* option.

The default forwarding set is used to forward packets of any FC in cases where all LSPs of the forwarding set the FC maps to become operationally DOWN. The router uses the user-configured default set as the initial default set. Otherwise, the router selects the lowest numbered set as the default forwarding set in a class-forwarding policy. When the last LSP in a default forwarding set goes into an operationally DOWN state, the router designates the next lowest-numbered set as the new default forwarding set.

A mapping to a class-forwarding policy and a set is added to the existing CBF configuration of an RSVP-TE or SR-TE LSP or to an LSP template. Use the following commands to perform this function:

- **MD-CLI**

```

configure router mpls lsp class-forwarding forwarding-set policy
configure router mpls lsp class-forwarding forwarding-set set
configure router mpls lsp-template class-forwarding forwarding-set policy
configure router mpls lsp-template class-forwarding forwarding-set set

```

- **classic CLI**

```

configure router mpls lsp class-forwarding forwarding-set policy set
configure router mpls lsp-template class-forwarding forwarding-set policy set

```

An MPLS LSP only maps to a single class-forwarding policy and forwarding set. Multiple LSPs can map to the same policy and set. If they form an ECMP set, from the IGP shortcut perspective, packets of the FCs mapped to this set are sprayed over these LSPs based on a modulo operation of the output of the hash routine on the headers of the packet and the number of LSPs in the set.

3.2.25.2 Feature behavior

When a VPN-v4/v6 prefix is resolved, the default behavior of the data path is to spray the packets over the entire ECMP set using a modulo operation of the number of resolved next hops in the ECMP set and the output of the hash on the packet header fields. With class-based forwarding enabled, the FC of the packet, is used to look up the forwarding set ID. Then, a modulo operation is performed on the tunnel next hops

of this set ID only, to spray packets of this FC. The data path concurrently implements ECMP within the tunnels of each set ID.

The CBF information of the LSPs forming the ECMP set is checked for consistency before programming. If more than a single class-forwarding policy exists, the set is considered inconsistent from a CBF perspective and no CBF information is programmed in the data-path and regular ECMP occurs.

Also, regardless of the CBF consistency check, the system programs the data-path with the full ECMP set.

The following describes the fallback behavior in data path of the CBF feature.

An FC, for which all LSPs in the forwarding set are operationally DOWN, has its packets forwarded over the default forwarding set. The default forwarding set is either the initial default forwarding set configured by the user or the lowest numbered set in the class-forwarding policy that has one or more LSPs in the operationally UP state. If the initial or subsequently elected default forwarding set has all its LSPs operationally DOWN, the next lower numbered forwarding set, which has at least one LSP in the operationally UP state, is elected as the default forwarding set.

If all LSPs of all forwarding sets become operationally DOWN, the router resumes regular ECMP spraying on the remaining LSPs in the full ECMP set.

Whenever the first LSP in a forwarding set becomes operationally UP, the router triggers the re-election of the default set and selects this set as the new default set, if it is the initial default set, otherwise, it selects the lowest numbered set.

SR OS implements a hierarchical ECMP architecture for BGP prefixes. The first level is the ECMP at the VPRN Level between different BGP next hop, and the second level is ECMP at the auto-bind-tunnel level, having the same next hop. This CBF feature is applied at the auto-bind-tunnel level. Weighted ECMP and the CBF feature are mutually exclusive on a per-BGP next-hop basis. When both are configured, Weighted ECMP takes the preference. CPM-originated packets on the router, including control plane and OAM packets, are forwarded over a single LSP from the set of LSPs that the packet's FC is mapped to, as per the CBF configuration.

3.2.26 IP VPN independent domains using BGP attribute set

In most IP VPN deployments, the network behind each CE router is relatively simple. In these networks, BGP is typically used to create an EBGP session between the CE device and service provider PE router. Occasionally, the service provider is required to provide an IP VPN service to a larger enterprise customer, or to another service provider that has complex sites using BGP for intra-site routing. Under such circumstances, it is useful to isolate the customer BGP domain from the service provider BGP domain, to ensure that routing policies applied in one domain do not affect the other domain. This functionality is achieved using BGP independent domains (based on RFC 6368), which introduce an optional transitive BGP path attribute called attribute set (ATTR_SET).

ATTR_SET offers benefits for both the service provider and customer. For example, ATTR_SET can hide the global AS number of the service provider from customer domains, even in inter-AS VPN scenarios. For the customer, ATTR_SET ensures that BGP routing decisions based on LOCAL_PREF or MED attributes (for example) are not affected when the service provider manipulates the same attributes in the core domain.

The following is the expected flow of route advertisements when a VPRN supports an independent domain:

1. The customer CE router advertises a BGP route to the service provider PE router, and it is received by the VPRN BGP instance. Although it is not expected, the BGP route may have an ATTR_SET attached. Use the following command to configure the router to remove ATTR_SETs from BGP routes received by the VPRN BGP instance:

- **MD-CLI**

```
configure service vprn bgp attribute-set remove true
```

- **classic CLI**

```
configure service vprn bgp attribute-set remove
```



Note: If the configuration of the **remove** command is changed, ROUTE_REFRESH messages are sent to all PE-CE peers of the VPRN.

2. If the BGP route in step 1 is matched and accepted by the VRF export policy of the VPRN (or the equivalent VRF target configuration takes effect), an ATTR_SET is added to the VPN-IP route created by the export process. Use the following command to configure the router to add ATTR_SETs to exported VPN-IP routes:

- **MD-CLI**

```
configure service vprn bgp-ipvpn attribute-set export true
```

- **classic CLI**

```
configure service vprn bgp-ipvpn attribute-set export
```

The ATTR_SET contains an exact copy of the BGP path attributes (post import policy) of the BGP route from the step 1, excluding the NEXT_HOP, MP_REACH, MP_UNREACH, AS4_PATH, and AS4_AGGREGATOR attributes. After the ATTR_SET is added to the VPN-IP route, the other path attributes of the VPN-IP route are initialized to the basic values that apply to exported local routes. The regenerated path attributes are influenced by the VRF export policy of the VPRN or the export policy that applies to the Base router BGP session carrying the VPN-IP routes, provided that the **vpn-apply-export** command is configured. Neither the VRF export policy nor a regular BGP export policy can modify the contents of the ATTR_SET.

3. The VPN-IP route in step 2 is received by another PE router and imported into a VPRN service that participates in the independent domain. Use the following command option to configure the VPRN service to accept and process ATTR_SETs in received VPN-IP routes.

```
configure service vprn bgp-ipvpn attribute-set import accept
```



Note: For a VPRN service not participating in an independent domain, Nokia recommends configuring the **import** command to **drop**. If the **import** command is configured to **accept**, only the attributes contained in the ATTR_SET influence the comparison of the route containing the ATTR_SET with other BGP routes in the VPRN BGP context. If BGP must compare an imported VPN-IP route containing an ATTR_SET to an imported VPN-IP route without an ATTR_SET, BGP compares ATTR_SET attributes against non-ATTR_SET attributes. However, Nokia recommends avoiding this scenario.

4. If the imported VPN-IP route in step 3 is the best overall route for the prefix, it is advertised to BGP CE peers of the VPRN. The following cases are exceptions where the best route with an ATTR_SET is not advertised to BGP CE peers:

- If the AS number of the VPRN is equal to the origin AS signaled inside the ATTR_SET, BGP routes with attributes derived from ATTR_SET are not advertised to non-client IBGP peers of the VPRN (peers not covered by a cluster configuration).
 - BGP routes with attributes derived from ATTR_SET are not advertised to confederation EBGP or IBGP peers of the VPRN.
5. In the routes advertised to BGP CE peers of the VPRN, the signaled attribute values are generally copies of the attribute values contained inside the ATTR_SET. However, the BGP export policy can modify the final values. If a BGP CE route is derived from a VPN-IP route with an ATTR_SET, the attributes in the route advertised to the CE are not based on the path attributes of the VPN-IP route.



Note: As per RFC 6368, when the AS number of the importing VPRN is not equal to the origin AS signaled inside the ATTR_SET, the origin AS is prepended to the AS path before advertising the route to the CE.

3.2.27 DHCP

This section describes DHCP feature support for VPRN services.

3.2.27.1 DHCP client for VPRN

The 7705 SAR Gen 2 supports VPRN interfaces configured with a DHCP client. When the node operates as a DHCP client, it learns the IP address of the interface via dynamic IP address assignment.

The DHCP client implementation for VPRN is identical to that of the Base router context. See "DHCP client" in the *7705 SAR Gen 2 Router Configuration Guide* for more information.

3.2.27.2 DHCP relay for VPRN

The 7705 SAR Gen 2 supports the DHCP relay agent for VPRN interfaces. When the DHCP relay agent is enabled, broadcast DHCP request messages are intercepted and forwarded unicast to a configured DHCP server.

The DHCP relay agent implementation for VPRN is identical to that of the Base router context. See "DHCP relay" in the *7705 SAR Gen 2 Router Configuration Guide* for more information.

3.2.27.3 Local DHCP server for VPRN

The 7705 SAR Gen 2 supports a local DHCP server for VPRN interfaces. When the DHCP relay agent is enabled, broadcast DHCP request messages are intercepted and can be forwarded unicast to a locally configured DHCP server.

See "DHCP relay" in the *7705 SAR Gen 2 Router Configuration Guide* for more information about the local DHCP server function.

See "DHCP persistence" in the *7705 SAR Gen 2 Basic System Configuration Guide* for more information about DHCP server persistency.

3.3 QoS on ingress bindings

Traffic is tunneled between VPRN service instances on different PEs over service tunnels bound to MPLS LSPs or GRE tunnels.

Use the commands in the following context to bind automatically the service tunnel to the underlying transport.

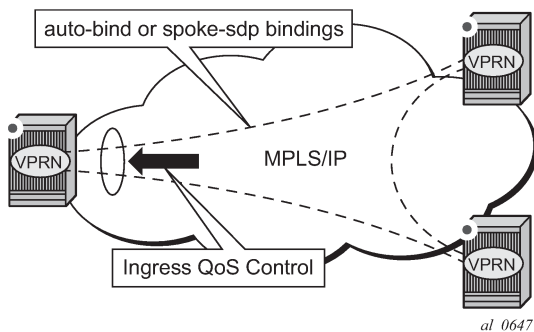
```
configure service vprn bgp-ipvpn mpls auto-bind-tunnel
```

Use the following command to bind statically the service tunnel to the underlying transport.

```
configure service vprn spoke-sdp
```

QoS control can be applied to the service tunnels for traffic ingressing into a VPRN service; see [Figure 27: Ingress QoS control on VPRN bindings](#).

Figure 27: Ingress QoS control on VPRN bindings



An ingress queue group must be configured and applied to the ingress network FP where the traffic is received for the VPRN. All traffic received on that FP for any binding in the VPRN (either automatically or statically configured) which is redirected to a policer in the FP queue group is controlled by that policer. Use the following command to configure the redirection in the network QoS policy.

```
configure qos network ingress fc fp-redirect-group
```

As a result, the traffic from all such bindings is treated as a single entity (per forwarding class) with regard to ingress QoS control. The following commands in the network QoS policy are ignored for this traffic (IP multicast traffic would use the ingress network queues or queue group related to the network interface).

```
configure qos network ingress fc fp-redirect-group broadcast-policer
configure qos network ingress fc fp-redirect-group mcast-policer
configure qos network ingress fc fp-redirect-group unknown-policer
```

Ingress classification is based on the configuration of the ingress section of the specified network QoS policy. The dot1p and exp classification is based on the outer Ethernet header and the MPLS label. The DSCP applies to the outer IP header if the tunnel encapsulation is GRE, or the DSCP in the first IP header

in the payload if you use the following command in the ingress section of the referenced network QoS policy.

```
configure qos network ingress ler-use-dscp
```

Ingress bandwidth control does not take into account the outer Ethernet header, the MPLS labels/control word or GRE headers, or the FCS of the incoming frame.

Use the following command to associate the network QoS policy and the FP queue group and instance within the network ingress of a VPRN.

```
configure service vprn network ingress qos fp-redirect-group instance
```

The preceding command overrides the QoS applied to the related network interfaces for unicast traffic arriving on bindings in that VPRN. The IP and IPv6 criteria statements are not supported in the applied network QoS policy.

This is supported for all available transport tunnel types and is independent of the allocation mode for VPRN service labels (**vrf** or **next-hop**) used within the VPRN. It is also supported for Carrier-Supporting-Carrier VPRNs.

The ingress network interfaces on which the traffic is received must be on FP2- and higher-based hardware.

3.4 Multicast in IP-VPN applications

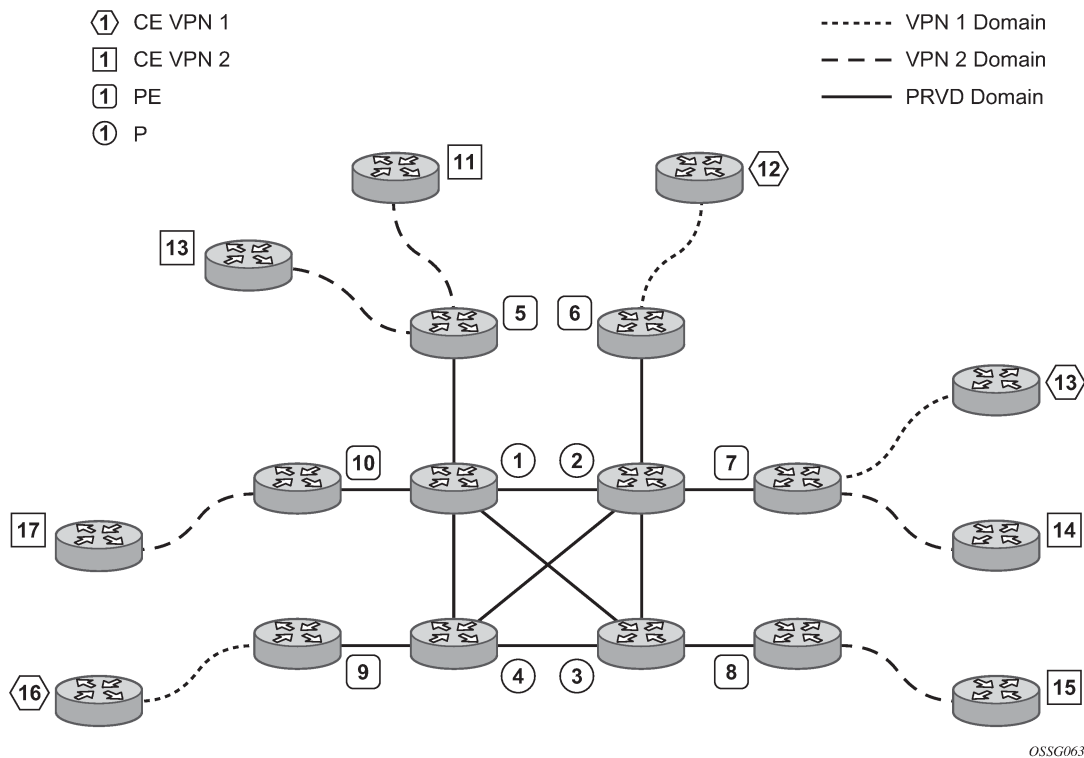
This section describes multicast in IP-VPN functionality. See the *7705 SAR Gen 2 Multicast Routing Protocols Guide* for more information about multicast protocols.

This feature can be used in the following types of applications: enterprise customers implementing a VPRN solution for their WAN networking needs; customer applications that include stock-ticker information; financial institutions for stock and other types of trading data; and video delivery systems.

Implementation of multicast in IP-VPNs requires the support and separation of the provider core multicast domain from customer multicast domains and the customer multicast domains from each other.

The following figure shows an example of multicast in an IP-VPN application.

Figure 28: Multicast in IP-VPN applications



The preceding figure shows the following domains:

- provider domain
 - core routers (1 through 4)
 - edge routers (5 through 10)
- customer IP-VPNs, each having its own multicast domain
 - VPN-1 (CE routers 12, 13, and 16)
 - VPN-2 (CE Routers 11, 14, 15, 17, and 18)

In this VPRN example, the VPN-1 data generated by the customer behind router 16 is multicast only by PE 9 to PE routers 6 and 7 for delivery to CE routers 12 and 13, respectively. VPN-2 data generated by the customer behind router 15 is forwarded by PE router 8 to PE routers 5, 7, and 10 for delivery to CE routers 18, 11, 14, and 17 respectively.

The demarcation points for these domains are in the PEs (routers 5 through 10). The PE routers participate in both the customer multicast domain and the provider multicast domain. The customer CEs are limited to a multicast adjacency with the multicast instance on the PE, where the PE multicast instance is specifically created to support that customer IP-VPN. As a result, customers are isolated from the provider core multicast domain and other customer multicast domains, while the provider core routers only participate in the provider multicast domain and are isolated from all customer multicast domains.

The PE for a customer multicast domain becomes adjacent to the CE routers attached to that PE and to all other PEs that participate in the IP-VPN (customer) multicast domain. The adjacencies are set up by the PE that encapsulates the customer multicast control data and the multicast streams inside the provider multicast packets. The encapsulated packets are forwarded only to the PE nodes that are attached to the

same customer edge routers as the originating stream and are part of the same customer VPRN. This process prunes the distribution of the multicast control and data traffic to the PEs that participate in the customer multicast domain.

The Rosen draft refers to this as the default multicast domain for this multicast domain; the multicast domain is associated with a unique multicast group address within the provider's network.

3.4.1 Use of data MDTs

Using the above method, all multicast data offered by a specific CE is always delivered to all other CEs that are part of the same multicast. It is possible that a number of CEs do not require the delivery of a particular multicast stream because they have no downstream receivers for a specific multicast group. At low traffic volumes, the impact of this is limited. However, at high data rates this could be optimized by devising a mechanism to prune PEs from the distribution tree that although forming part of the customer multicast have no need to deliver a specific multicast stream to the CE attached to them. To facilitate this optimization, the Rosen draft specifies the use of data MDTs. These data MDTs are signaled after the bandwidth for a specific SG exceeds the configurable threshold.

When a PE detects it is transmitting data for the SG in excess of this threshold, it sends an MDT join TLV (at 60 second intervals) over the default MDT to all PEs. All PEs that require the SG specified in the MDT join TLV join the data MDT that is used by the transmitting PE to send the specific SG. PEs that do not require the SG do not join the data MDT, therefore pruning the multicast distribution tree to just the PEs requiring the SG. After providing sufficient time for all PEs to join the data MDT, the transmitting PE switches the specific multicast stream to the data MDT.

PEs that do not require the SG to be delivered, keep state to allow them to join the data MDT as required.

When the bandwidth requirement no longer exceeds the threshold, the PE stops announcing the MDT join TLV. At this point the PEs using the data MDT leave this group and transmission resumes over the default MDT.

Sampling to check if an s,g has exceeded the threshold occurs every ten seconds. If the rate has exceeded the configured rate in that sample period then the data MDT is created. If during that period the transmission rate has not exceeded the configured threshold then the data MDT is not created. If the data MDT is active and the transmission rate in the last sample period has not exceeded the configured rate then the data MDT is torn down and the multicast stream resumes transmission over the default MDT.

3.4.2 Multicast protocols supported in the provider network

When MVPN auto-discovery is disabled, PIM-SM can be used for I-PMSI, and PIM-SSM or PIM-SM (Draft-Rosen Data MDT) can be used for S-PMSI; When MVPN S-PMSI auto-discovery is enabled, both PIM-SM and PIM SSM can be used for I-PMSI, and PIM-SSM can be used for S-PMSI. In the customer network, both PIM-SM and PIM-SSM are supported.

An MVPN is defined by two sets of sites: sender sites set and receiver sites set, with the following properties:

- Hosts within the sender sites set could originate multicast traffic for receivers in the receiver sites set.
- Receivers not in the receiver sites set should not be able to receive this traffic.
- Hosts within the receiver sites set could receive multicast traffic originated by any host in the sender sites set.

- Hosts within the receiver sites set should not be able to receive multicast traffic originated by any host that is not in the sender sites set.

A site could be both in the sender sites set and receiver sites set, which implies that hosts within such a site could both originate and receive multicast traffic. An extreme case is when the sender sites set is the same as the receiver sites set, in which case all sites could originate and receive multicast traffic from each other.

Sites within a specific MVPN may be either within the same, or in different organizations, which implies that an MVPN can be either an intranet or an extranet. A site may be in more than one MVPN, which implies that MVPNs may overlap. Not all sites of a specific MVPN have to be connected to the same service provider, which implies that an MVPN can span multiple service providers.

Another way to look at MVPN is to say that an MVPN is defined by a set of administrative policies. Such policies determine both sender sites set and receiver site set. Such policies are established by MVPN customers, but implemented by MVPN service providers using the existing BGP/MPLS VPN mechanisms, such as route targets, with extensions, as necessary.

3.4.3 MVPN membership autodiscovery using BGP

BGP-based autodiscovery is performed by a multicast VPN address family. Any PE that attaches to an MVPN must issue a BGP update message containing an NLRI in this address family, along with a specific set of attributes.

The PE router uses route targets to specify MVPN route import and export. The route target may be the same as the one used for the corresponding unicast VPN, or it may be different. The PE router can specify separate import route targets for sender sites and receiver sites for a specific MVPN.

The route distinguisher (RD) that is used for the corresponding unicast VPN can also be used for the MVPN.

When BGP autodiscovery is enabled, PIM peering on the I-PMSI is disabled, so no PIM hellos are sent on the I-PMSI. C-trees to P-tunnels bindings are also discovered using BGP S-PMSI AD routes, instead of PIM join TLVs.

SR OS NG-MVPN supports the autodiscovery default (BGP) or MDT-SAFI (PIM). Use the following commands to configure these options.

```
configure service vprn mvpn auto-discovery
```

```
configure service vprn mvpn c-mcast-signaling pim
```

[Table 6: Supported configuration combinations](#) and [Table 7: Supported configuration combinations](#) describe the supported configuration combinations. If the CLI combination is not allowed, the system returns an error message. If the CLI command is marked as "ignored" in the table, the configuration is not blocked, but its value is ignored by the software.

Table 6: Supported configuration combinations

Auto-discovery	Inclusive PIM SSM	Action
Yes	Yes	Allowed
MDT-SAFI	Yes	Allowed

Auto-discovery	Inclusive PIM SSM	Action
No	Yes	Not Allowed
Yes or No	No	Allowed
MDT-SAFI	No	Ignored
MDT-SAFI	No (RSVP and MLDP)	Not Allowed

Table 7: Supported configuration combinations

Auto-discovery	C-mcast-signaling	s-PMSI auto-discovery	Action
Yes	BGP	Ignored	Allowed
Yes	PIM	Yes	Allowed
Yes	PIM	No	Allowed
No	BGP	Ignored	Not Allowed
No	PIM	Ignored	Allowed
MDT-SAFI	Ignored (PIM behavior)	Ignored ("No" behavior)	Allowed

For example, if autodiscovery is disabled, the **c-mcast-signaling bgp** command fails with an error message stating:

C-multicast signaling in BGP requires autodiscovery to be enabled.

If **c-mcast-signaling bgp** is configured, disabling autodiscovery in the following context fails.

```
configure service vprn mvpn provider-tunnel selective
```

The error message states as follows.

C-multicast signaling in BGP requires autodiscovery to be enabled

When **c-mcast-signaling bgp** is configured, S-PMSI A-D is always enabled (its configuration is ignored).

When autodiscovery is disabled, S-PMSI A-D is always disabled (its configuration is ignored).

When autodiscovery is enabled and **c-multicast-signaling pim** is configured, the S-PMSI A-D configuration value is used.

MDT-SAFI uses PIM C-multicast signaling and S-PMSI signaling regardless of what is configured. A C-multicast signaling or S-PMSI signaling configuration is ignored, but both **pim** and **bgp** command options are allowed.

MDT-SAFI is only applicable to PIM-SSM I-PMSI. PIM-SM (ASM) I-PMSI is configurable but is ignored. RSVP and MLDP I-PMSI are not allowed.

MVPN implementation based on RFC 6037, *Cisco Systems' Solution for Multicast in MPLS/BGP IP VPNs* can support membership autodiscovery using BGP MDT-SAFI. A CLI option is provided per MVPN instance to enable auto-discovery using either BGP MDT-SAFI or NG-MVPN. Only PIM-MDT is supported with the BGP MDT-SAFI method.

3.4.4 PE-PE transmission of C-multicast routing using BGP

MVPN C-multicast routing information is exchanged between PEs by using C-multicast routes that are carried using MCAST-VPN NLRI.

3.4.5 VRF route import extended community

VRF route import is an IP address-specific extended community (extended type) and is transitive across AS boundaries, as described in RFC 4360, *BGP Extended Communities Attribute*.

To support MVPN, in addition to the import and export route target extended communities used by unicast routing, each VRF on a PE must have an import route target extended community that controls the import of C-multicast routes into a VRF.

The C-multicast import routing table (RT) uniquely identifies a VRF and is set up as follows:

- The Global Administrator field of the C-multicast import RT must be set to an IP address of the PE. This address should be common for all the VRFs on the PE (this address may be the PE loopback address).
- The Local Administrator field of the C-multicast import RT associated with a VRF contains a two-octet number that uniquely identifies that VRF within the PE that contains the VRF.

A PE that has sites of an MVPN connected to it communicates the value of the C-multicast import RT associated with the VRF of that MVPN on the PE to all other PEs that have sites on that MVPN.

A PE that originates a (unicast) route to VPN IP addresses includes (in the BGP update message that carries this route) the VRF route import extended community that has the value of the C-multicast import RT of the VRF associated with the route. However, if none of these addresses act as multicast sources or an RP, the (unicast) route does not need to carry the VRF route import extended community.

All C-multicast routes with the C-multicast import RT specific to the VRF must be accepted. VRF import and VRF target policies do not apply to C-multicast routes.

The following example shows the decision flow path.

```
if (route-type == c-mcast-route)
  if (route_target_list includes C-multicast_Import_RT){
    else
      drop;
    else
      Run vrf_import and/or vrf-target;
```

3.4.6 Provider tunnel support

3.4.6.1 Point-to-Multipoint Inclusive (I-PMSI) and Selective (S-PMSI) Provider Multicast Service Interface

BGP C-multicast signaling must be enabled for an MVPN instance to use P2MP RSVP-TE or LDP as I-PMSI (equivalent to 'Default MDT', as defined in draft Rosen MVPN) and S-PMSI (equivalent to 'Data MDT', as defined in draft Rosen MVPN).

By default, all PE nodes participating in an MVPN receive data traffic over I-PMSI. Optionally, (C-*, C-*) wildcard S-PMSI can be used instead of I-PMSI. See section [Wildcard \(C-*, C-*\) P2MP LSP S-PMSI](#) for more information. For efficient data traffic distribution, one or more S-PMSIs can be used, in addition to the default PMSI, to send traffic to PE nodes that have at least one active receiver connected to them. For more information, see [P2MP LSP S-PMSI](#).

Only one unique multicast flow is supported over each P2MP RSVP-TE or P2MP LDP LSP S-PMSI. Number of S-PMSI that can be initiated per MVPN instance is restricted by the **maximum-p2mp-spmsi** command. P2MP LSP S-PMSI cannot be used for more than one (S,G) stream (that is, multiple multicast flow) as number of S-PMSI per MVPN limit is reached. Multicast flows that cannot switch to S-PMSI remain on I-PMSI.

3.4.6.2 P2MP RSVP-TE I-PMSI and S-PMSI

Point-to-Multipoint RSVP-TE LSP as inclusive or selective provider tunnel is available with BGP NG-MVPN only. P2MP RSVP-TE LSP is dynamically setup from root node on auto discovery of leaf PE nodes that are participating in multicast VPN. Each RSVP-TE I-PMSI or S-PMSI LSP can be used with a single MVPN instance only.

RSVP-TE LSP template must be defined (see *7705 SAR Gen 2 MPLS Guide*) and bound to MVPN as inclusive or selective (S-PMSI is for efficient data distribution and is optional) provider tunnel to dynamically initiate P2MP LSP to the leaf PE nodes learned via NG-MVPN auto-discovery signaling. Each P2MP LSP S2L is signaled based on parameters defined in LSP template.

3.4.6.3 P2MP LDP I-PMSI and S-PMSI

P2MP LDP LSP as an inclusive or selective provider tunnel is available with BGP NG-MVPN only. P2MP LDP LSP is dynamically set up from leaf nodes after the auto-discovery of leaf node PE nodes that are participating in MVPN. Each LDP I-PMSI or S-PMSI LSP can be used with only a single MVPN instance.

P2MP LDP must be configured as an inclusive or selective (S-PMSI is for efficient data distribution and is optional) provider tunnel per MVPN to dynamically initiate P2MP LSP to leaf PE nodes learned via NG-MVPN auto-discovery signaling.

3.4.6.4 Wildcard (C-*, C-*) P2MP LSP S-PMSI

Wildcard S-PMSI allows the use of selective tunnel as a default tunnel for a specific MVPN. Users can avoid a full mesh of LSPs between the MVPN PEs, reducing related signaling, state, and bandwidth

consumption for multicast distribution. No traffic is sent to PEs unless receivers are active on the default PMSI.

Use the following command to configure the wildcard S-PMSI functionality for NG-MVPN using LDP and RSVP-TE in P-instance.

```
configure service vprn mvpn provider-tunnel inclusive wildcard-spmsi
```

The support includes:

- IPv4 and IPv6
- PIM ASM and SSM
- directly attached receivers

The (C-*, C-*) wildcard implementation uses wildcard S-PMSI instead of I-PMSI for a specific MVPN. A VPRN shutdown is required to switch MVPN from I-PMSI to (C-*, C-*) S-PMSI. ISSU and Upstream Multicast Hop (UMH) redundancy can be used to minimize the impact.

To minimize outage, the following upgrade order is recommended:

1. Route Reflector
2. receiver PEs
3. backup UMH
4. active UMH

Use the following command to configure RSVP-TE/mLDP under the inclusive provider tunnel. The configuration applies to the wildcard S-PMSI when enabled.

```
configure service vprn mvpn provider-tunnel inclusive
```

Wildcard C-S and C-G values are encoded as defined in RFC6625; that is, using zero for Multicast Source Length and Multicast Group Length, and omitting Multicast Source and Multicast Group values respectively in MCAST_VPN_NLRI. For example, a (C-*, C-*) is advertised as: RD, 0x00, 0x00, and the IP address of the originating router.

Procedures implemented by SR OS are compliant with section 3 and 4 of RFC 6625. Wildcards encoded as described in the preceding paragraph are carried in the NLRI field of MP_REACH_NLRF_ATTRIBUTE. Both IPv4 and IPv6 are supported: (AFI) of 1 or 2 and a Subsequent AFI (SAFI) of MCAST-VPN.

The (C-*, C-*) S-PMSI is established as follows:

- UMH PEs advertise I-PMSI A-D routes without tunnel information present (empty PTA), and encoded in accordance with RFC 6513 and RFC 6514, before advertising wildcard S-PMSI. I-PMSI must be signaled and installed on the receiver PEs because (C-*, C-*) S-PMSI is only installed when a first receiver is added. However, no LSP is established for I-PMSI.
- UMH PEs advertise the S-PMSI A-D route whose NLRI contains (C-*, C-*) with tunnel information encoded in accordance with RFC 6625.
- Receiver PEs join wildcard S-PMSI if receivers are present.



Note: If the UMH PE does not encode I-PMSI/S-PMSI A-D routes as described, or advertises both the I-PMSI and wildcard S-PMSI with the tunnel information present, the interoperability cannot be achieved.

To ensure correct operation between PEs with (C-*, C-*) S-PMSI signaling, two BSR modes of operation are implemented. These modes are the BSR unicast (default) and BSR S-PMSI.

The following applies to the BSR unicast mode:

- BSR PDUs are sent or forwarded as unicast PDUs to neighbor PEs when I-PMSI with pseudo-tunnel interface is installed.
- At every BSR interval timer, the BSR Unicast PDUs are sent to all I-PMSI interfaces when this is an elected BSR.
- BSMs received as multicast from C-instance interfaces are flooded as unicast in the P-instance.
- All PEs process BSR PDUs received on the I-PMSI pseudo-tunnel interface as unicast packets.
- BSR PDUs are not forwarded to the management control interface of the PE.
- BSR unicast PDUs use PE's system IP address of the PE as the destination IP, and sender the system address of the sender PE as the source IP.
- The BSR unicast functionality ensures that no special state needs to be created for BSR when (C-*, C-*) S-PMSI is enabled, which is beneficial considering the low volume of BSR traffic.



Note:

- For BSR unicast, the base IPv4 system address (IPv4) or the mapped version of the base IPv4 system address (IPv6) must be configured under the VPRN to ensure BSR unicast messages can reach the VPRN.
- For BSR S-PMSI, the base IPv4 or IPv6 system address must be configured under the VPRN to ensure BSR S-PMSIs are established.

The BSR S-PMSI mode can be enabled to allow interoperability with other vendors. In this mode, full mesh S-PMSI is required and created between all PEs in MVPN to exchange BSR PDUs. To operate as expected, the BSR S-PMSI mode requires a selective P-tunnel configuration. For IPv6 support (including dual-stack) of BSR S-PMSI mode, the IPv6 default system interface address must be configured as a loopback interface address under the VPRN and VPRN PIM contexts. Changing the BSR signaling requires a VPRN shutdown.

Other key feature interactions and restrictions for (C-*, C-*) include the following:

- Extranet is fully supported with wildcard S-PMSI trees.
- (C-S, C-G) S-PMSIs are supported when (C-*, C-*) S-PMSI is configured (including both BW and receiver PE driven thresholds).
- Geo-redundancy is supported (deploying with geo-redundancy eliminates traffic duplication when geo-redundant source has no active receivers at a cost of slightly increased outage upon a switch because wildcard S-PMSI may need to be re-establish).
- PIM in P-instance is not supported.
- The implementation requires wildcard encoding as described in RFC 6625 and I-PMSI/S-PMSI signaling as defined above (I-PMSI signaled with empty PTA then S-PMSI signaled with P-tunnel PTA) for interoperability. Implementations that do not adhere to the RFC 6625 encoding, or signal both I-PMSI and S-PMSI with P-tunnel PTA, do not interoperate with the SR OS implementation).

3.4.6.5 P2MP LSP S-PMSI

NG-MVPN support P2MP RSVP-TE and P2MP LDP LSPs as selective provider multicast service interface (S-PMSI). S-PMSI is used to avoid sending traffic to PEs that participate in multicast VPN, but do not have any receivers for a specific C-multicast flow. This allows more-BW efficient distribution of multicast traffic over the provider network, especially for high bandwidth multicast flows. S-PMSI is spawned dynamically based on configured triggers as described in [S-PMSI trigger thresholds](#).

In MVPN, the head-end PE discovers all the leaf PEs via I-PMSI A-D routes. It then signals the P2MP LSP to all the leaf PEs using RSVP-TE. In the scenario of S-PMSI:

1. The head-end PE sends an S-PMSI A-D route for a specific C-flow with the Leaf Information Required bit set.
2. The PEs interested in the C-flow respond with Leaf A-D routes.
3. The head-end PE signals the P2MP LSP to all the leaf PEs using RSVP-TE.

Because the receivers may come and go, the implementation supports dynamically adding and pruning leaf nodes to and from the P2MP LSP.

When the tunnel type in the PMSI attribute is set to RSVP-TE P2MP LSP, the tunnel identifier is <Extended Tunnel ID, Reserved, Tunnel ID, P2MP ID>, as carried in the RSVP-TE P2MP LSP SESSION Object.

The PE can also learn via an A-D route that it needs to receive traffic on a particular RSVP-TE P2MP LSP before the LSP is actually set up. In this case, the PE must wait until the LSP is operational before it can modify its forwarding tables as directed by the A-D route.

Because of the way that LDP normally works, mLDP P2MP LSPs are set up without solicitation from the leaf PEs toward the head-end PE. The leaf PE discovers the head-end PE via I-PMSI or S-PMSI A-D routes. The tunnel identifier carried in the PMSI attribute is used as the P2MP FEC element. The tunnel identifier consists of the address of the head-end PE and a Generic LSP identifier value. The Generic LSP identifier value is automatically generated by the head-end PE.

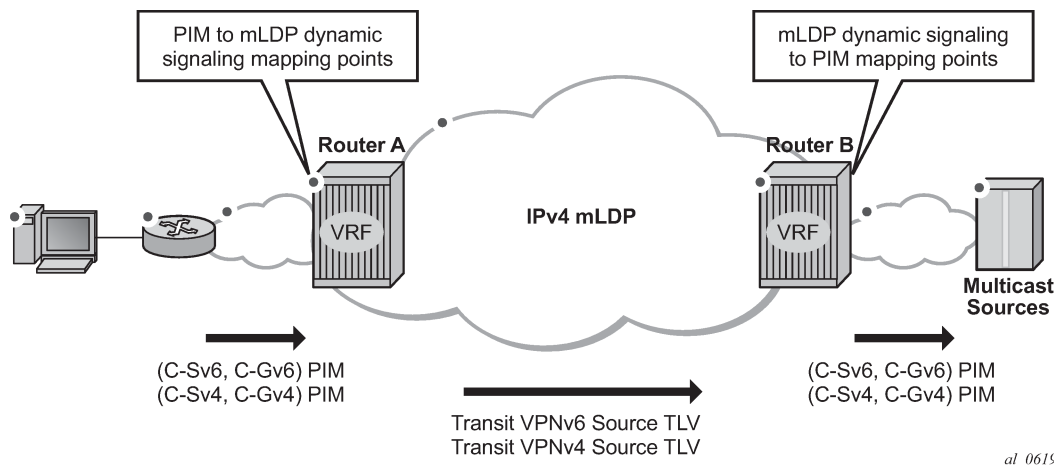
3.4.6.6 Dynamic multicast signaling over P2MP LDP in VRF

This feature provides a multicast signaling solution for IP-VPNs, allowing the connection of IP multicast sources and receivers in C-instances, which are running PIM multicast protocol using Rosen MVPN with BGP SAFI and P2MP mLDP in P-instance. The solution dynamically maps each PIM multicast flow to a P2MP LDP LSP on the source and receiver PEs.

The feature uses procedures defined in RFC 7246: *Multipoint Label Distribution Protocol In-Band Signaling in Virtual Routing and Forwarding (VRF) Table Context*. On the receiver PE, PIM signaling is dynamically mapped to the P2MP LDP tree setup. On the source PE, signaling is handed back from the P2MP mLDP to the PIM. Because of dynamic mapping of multicast IP flow to P2MP LSP, provisioning and maintenance overhead is eliminated as multicast destination services are added and removed from the VRF. Per (C-S, C-G) IP multicast state is also removed from the network, because P2MP LSPs are used to transport multicast flows.

[Figure 29: Dynamic mLDP signaling for IP multicast in VPRN](#) illustrates dynamic mLDP signaling for IP multicast in VPRN.

Figure 29: Dynamic mLDP signaling for IP multicast in VPRN



As illustrated in [Figure 29: Dynamic mLDP signaling for IP multicast in VPRN](#), P2MP LDP LSP signaling is initiated from the receiver PE that receives PIM JOIN from a downstream router (Router A). Use the commands in the following context to enable dynamic multicast signaling on PIM customer-facing interfaces for the specific VPRN of Router A.

```
configure service vprn pim interface p2mp-ldp-tree-join
```

This enables handover of multicast tree signaling from the PIM to the P2MP LDP LSP. Being a leaf node of the P2MP LDP LSP, Router A selects the upstream-hop as the root node of P2MP LDP FEC, based on a routing table lookup. If an ECMP path is available for a specific route, then the number of trees are equally balanced toward multiple root nodes. The PIM joins are carried in the Transit Source PE (Router B), and multicast tree signaling is handed back to the PIM and propagated upstream as native-IP PIM JOIN toward C-instance multicast source.

The feature is supported with IPv4 and IPv6 PIM SSM and IPv4 mLDP. Directly connected IGMP/MLD receivers are also supported, with PIM enabled on outgoing interfaces and SSM mapping configured, if required.

The following are feature restrictions:

- Dynamic mLDP signaling in a VPRN instance and Rosen or NG-MVPN are mutually exclusive.
- A single instance of P2MP LDP LSP is supported between the receiver PE and Source PE per multicast flow; there is no stitching of dynamic trees.
- Extranet functionality is not supported.
- The router LSA link ID or the advertising router ID must be a routable IPv4 address (including IPv6 into IPv4 mLDP use cases).
- IPv6 PIM with dynamic IPv4 mLDP signaling is not supported with EBGp or IBGP with IPv6 next-hop.
- Inter-AS and IGP inter-area scenarios where the originating router is altered at the ASBR and ABR respectively, (therefore PIM has no way to create the LDP LSP toward the source), are not supported.
- When dynamic mLDP signaling is deployed, a change in Route Distinguisher (RD) on the Source PE is not acted upon for any (C-S, C-G)s until the receiver PEs learn about the new RD (via BGP) and send explicit delete and create with the new RD.

- Procedures of Section 2 of RFC 7246 for a case where UMH and the upstream PE do not have the same IP address are not supported.

3.4.6.7 MVPN sender-only/receiver-only

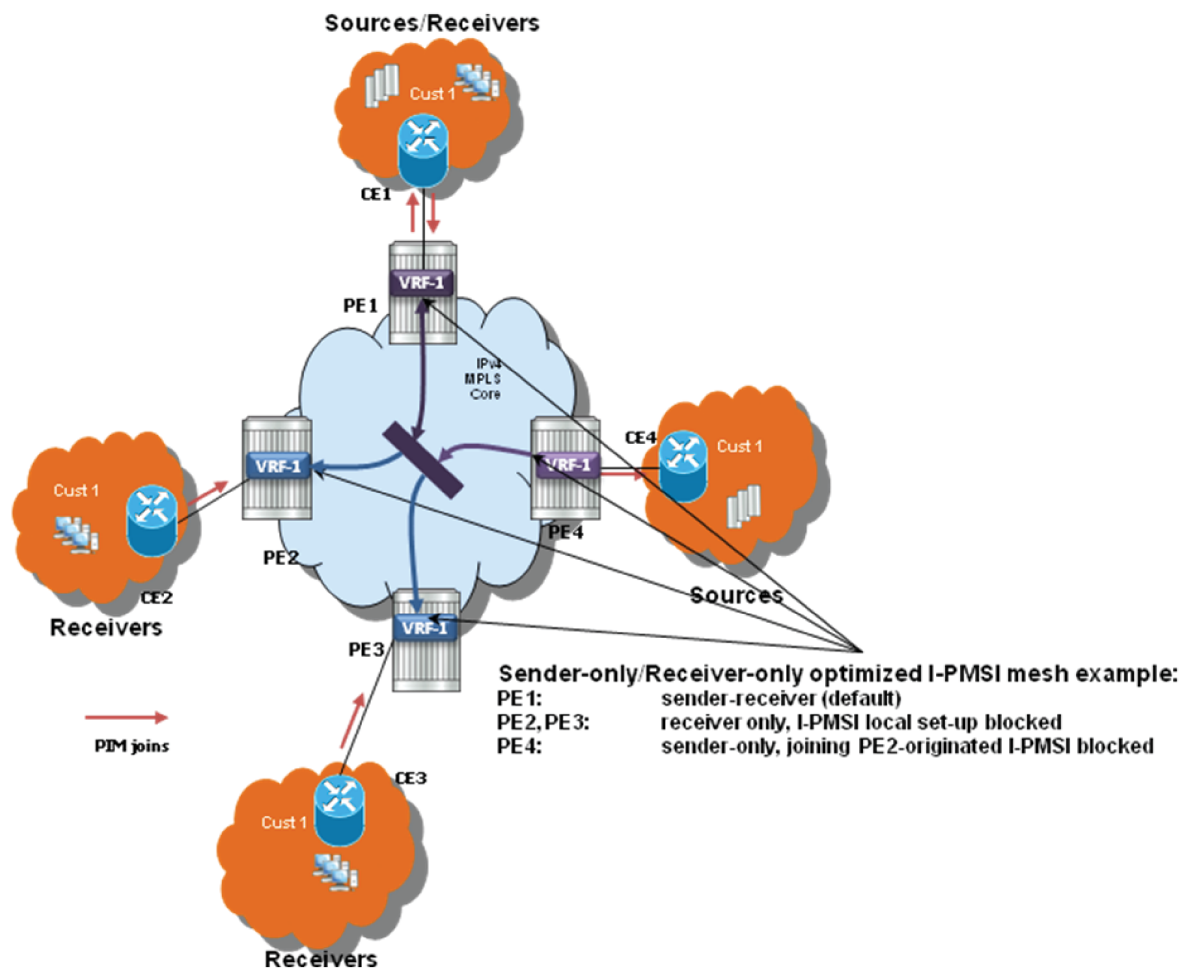
In multicast MVPN, if multiple PE nodes form a peering with a common MVPN instance, each PE node originates by default a multicast tree locally toward the remaining PE nodes that are members of this MVPN instance. This behavior creates a mesh of I-PMSI across all PE nodes in the MVPN. It is often the case that a specific VPN has many sites that host multicast receivers, but only a few sites that either host both receivers and sources or sources only.

MVPN Sender-only/Receiver-only allows optimization of control and data plane resources by preventing unnecessary I-PMSI mesh when a specific PE hosts multicast sources only or multicast receivers only for a specific MVPN.

For PE nodes that host only multicast sources for a specific VPN, the user can block those PEs, through configuration, from joining I-PMSIs from other PEs in this MVPN. For PE nodes that host only multicast receivers for a specific VPN, the user can block those PEs, through configuration, to set-up a local I-PMSI to other PEs in this MVPN.

MVPN Sender-only/Receiver-only is supported with NG-MVPN using IPv4 RSVP-TE or IPv4 LDP provider tunnels for both IPv4 and IPv6 customer multicast. [Figure 30: MVPN sender-only/receiver-only example](#) depicts 4-site MVPN with sender-only, receiver-only and sender-receiver (default) sites.

Figure 30: MVPN sender-only/receiver-only example



Extra attention needs to be paid to BSR/RP placement when Sender-only/Receiver-only is enabled. Source DR sends unicast encapsulated traffic toward RP, therefore, RP shall be at sender-receiver or sender-only site, so that *G traffic can be sent over the tunnel. BSR shall be deployed at the sender-receiver site. BSR can be at sender-only site if the RPs are at the same site. BSR needs to receive packets from other candidate-BSR and candidate-RPs and also needs to send BSM packets to everyone.

3.4.6.8 S-PMSI trigger thresholds

The mLDP and RSVP-TE S-PMSIs support two types of data thresholds: bandwidth-driven and receiver-PE-driven. The threshold evaluation and bandwidth driven threshold functionality are described in [Use of data MDTs](#).

In addition to the bandwidth threshold functionality, the user can enable receiver-PE-driven threshold behavior. Receiver PE thresholds ensure that S-PMSI is only created when BW savings in P-instance justify extra signaling required to establish a new S-PMSI. For example, the number of receiver PEs interested in a specific C-multicast flow is meaningfully smaller than the number of receiver PEs for default PMSI (I-PMSI or wildcard S-PMSI). To ensure that S-PMSI is not constantly created/deleted, two

thresholds need to be specified: receiver PE add threshold and receiver PE delete threshold (expected to be significantly higher).

When a (C-S, C-G) crosses a data threshold to create S-PMSI, instead of regular S-PMSI signaling, sender PE originates S-PMSI explicit tracking procedures to detect how many receiver PEs are interested in a specific (C-S, C-G). When receiver PEs receive an explicit tracking request, each receiver PE responds, indicating whether there are multicast receivers present for that (C-S, C-G) on the specific PE (PE is interested in a specific (C-S, C-G)). If the geo-redundancy feature is enabled, receiver PEs do not respond to explicit tracking requests for suppressed sources and therefore only Receiver PEs with an active join are counted against the configured thresholds on Source PEs.

Upon regular sampling and check interval, if the previous check interval had a non-zero receiver PE count (one interval delay to not trigger S-PMSI prematurely) and current count of receiver PEs interested in the specific (C-S, C-G) is non-zero and is less than the configured receiver PE add threshold, Source PE sets up S-PMSI for this (C-S, C-G) following standard ng-MVPN procedures augmented with explicit tracking for S-PMSI being established.

Data threshold timer should be set to ensure enough time is given for explicit tracking to complete (for example, setting the timer to value that is too low may create S-PMSI prematurely).

Upon regular data-delay-interval expiry processing, when BW threshold validity is being checked, a current receiver PE count is also checked (for example, explicit tracking continues on the established S-PMSI). If BW threshold no longer applies or the receiver PEs exceed receiver PE delete threshold, the S-PMSI is torn down and (C-S, C-G) joins back the default PMSI.

Changing of thresholds (including enabling/disabling the thresholds) is allowed in service. The configuration change is evaluated at the next periodic threshold evaluation.

The explicit tracking procedures follow RFC 6513/6514 with clarification and wildcard S-PMSI explicit tracking extensions as described in IETF Draft: *draft-dolganow-l3vpn-expl-track-00*.

3.4.6.9 Migration from existing Rosen implementation

The existing Rosen implementation is compatible to provide an easy migration path.

The following migration procedures are supported:

- Upgrade all the PE nodes that need to support MVPN to the newer release.
- The old configuration is converted automatically to the new style.
- Node by node, MCAST-VPN address-family for BGP is enabled. Enable auto-discovery using BGP.
- Change PE-to-PE signaling to BGP.

3.4.6.10 Policy-based S-PMSI

SR OS creates a single selective P-Multicast Service Interface (S-PMSI) per multicast stream: (S,G) or (*,G). To better manage bandwidth allocation in the network, multiple multicast streams are often bundled starting from the same root node into a single, multi-stream S-PMSI. Network bandwidth is usually managed per package or group of packages, instead of per channel.

Multi-stream S-PMSI supports a single S-PMSI for one or more IPv4 (C-S, C-G) or IPv6 (C-S, C-G) stream. Multiple multicast streams starting from the same VPRN going to the same set of leafs can use a single S-PMSI. Multi-stream S-PMSIs can:

- carry exclusively IPv4 or exclusively IPv6, or a mix of channels

- coexist with a single group S-PMSI

To create a multi-stream S-PMSI, an S-PMSI policy needs to be configured in the VPN context on the source node. This policy maps multiple (C-S, C-G) streams to a single S-PMSI. Because this configuration is done per MVPN, multiple VPNs can have identical policies, each configured for its own VPN context.

When mapping a multicast stream to a multi-stream S-PMSI policy, the data traverses the S-PMSI without first using the I-PMSI. (Before this feature, when a multicast stream was sourced, the data used the I-PMSI first until a configured threshold was met. After this, if the multicast data exceeded the threshold, it signaled the S-PMSI tunnel and switched from I-PMSI to S-PMSI.)

For multi-stream S-PMSI, if the policy is configured and the multicast data matches the policy rules, the multi-stream S-PMSI is used from the start without using the default I-PMSI.

Multiple multi-stream S-PMSI policies could be assigned to a specific S-PMSI configuration. In this case, the policy acts as a link list. The first (lowest index) that matches the multi-stream S-PMSI policy is used for that specific stream.

The rules for matching a multi-stream S-PMSI on the source node are listed in this section.

S-PMSI to (C-S, C-G) mapping on Source-PE, in sequence:

1. The multi-stream S-PMSI policy is evaluated, starting from the lowest numerical policy index. This allows the feature to be enabled in the service when per-(C-S, C-G) stream configuration is present. Only entries that are not shut down are evaluated. First, the multi-stream S-PMSI (the lowest policy index) that the (C-S, C-G) stream maps to is selected.
2. If (C-S, C-G) does not map to any of the multi-stream S-PMSIs, per-(C-S, C-G) S-PMSIs are used for transmission if a (C-S, C-G) maps to an existing S-PMSI (based on data-thresholds).
3. If S-PMSI cannot be used, the default PMSI is used.

Multi-stream S-PMSI P-tunnel failure handling can be performed. If an S-PMSI P-tunnel is not available, a default PMSI tunnel is used. When an S-PMSI tunnel fails, all (C-S, C-G) streams using this multi-stream S-PMSI move to the default PMSI. The groups move back to S-PMSI after the S-PMSI tunnel is restored.

3.4.6.10.1 Supported MPLS tunnels

Multi-stream S-PMSI is configured in the context of an auto-discovery default (that is, NG-MVPN). It supports all existing per-mLDP/RSVP-TE P2MP S-PMSI tunnel functionality for multi-stream S-PMSI LSP templates (RSVP-TE P-instance).



Note:

- Per-multicast group statistics are not available for multi-stream S-PMSIs on S-PMSI level.
- GRE tunnels are not supported for multi-stream S-PMSI.

3.4.6.10.2 Supported multicast features

S-PMSI is supported with PIM-ASM and PIM-SSM in C-instances.



Note:

- The multi-stream S-PMSI model uses BSR RP co-located with the source PE or an RP between the source PE and multicast source, that is, upstream of receivers. Both **bsr unicast** and **bsr spmsi** can be deployed as applicable.

- The model also supports other RP types.

3.4.6.10.3 In-service changes to multi-stream S-PMSI

The operator can change the mapping in service by moving active streams (C-S, C-G) from one S-PMSI to another using the configuration, or from the default PMSI to the S-PMSI, without having to stop data transmission or having to disable a PMSI.

The change is performed by moving a (C-S, C-G) stream from a per-group S-PMSI to a multi-stream S-PMSI and the reverse, and moving a (C-S, C-G) stream from one multi-stream S-PMSI to another multi-stream S-PMSI.



Note:

- During re-mapping, a changed (C-S, C-G) stream is first moved to the default PMSI before it is moved to a new S-PMSI, regardless of the type of move. Unchanged (C-S, C-G) streams must remain on an existing PMSIs.
- Any change to a multi-stream S-PMSI policy or to a preferred multi-stream S-PMSI policy (for example, an index change equivalent to less than or equal to the current policy) should be performed during a maintenance window. Failure to perform these types of changes in a maintenance window could potentially cause a traffic outage.

3.4.6.10.4 Configuration example

In the following example, two policies are created on the source node: multi-stream S-PMSI 1 and multi-stream-S-PMSI 10.

A multicast stream with group 224.0.0.0 and source 192.0.2.0/24 maps to the first multi-stream policy. The group in the range of 224.0.0.0/24 and source 192.0.1.0/24 maps to policy 10.

Example: MD-CLI

```
[ex:/configure service vprn "5" mvpn]
A:admin@node-2# info
  c-mcast-signaling bgp
  auto-discovery {
    type bgp
  }
  provider-tunnel {
    inclusive {
      mldp {
        admin-state enable
      }
    }
  }
  selective {
    auto-discovery true
    data-threshold {
      group-prefix 239.0.0.0/8 {
        threshold 1
      }
      group-prefix 239.70.1.0/24 {
        threshold 1
      }
    }
  }
  multistream-spmsi 1 {
    admin-state enable
  }
```

```

        group-prefix 224.0.0.0/24 source-prefix 192.0.2.0/24 { }
    }
    multistream-spmsi 10 {
        admin-state enable
        group-prefix 224.0.0.0/24 source-prefix 192.0.1.0/24 { }
    }
    mldp {
        admin-state enable
    }
}

```

Example: classic CLI

```

A:node-2:>config>service>vprn>mvpn# info
    auto-discovery default
    c-mcast-signaling bgp
    provider-tunnel
        inclusive
        mldp
            no shutdown
        exit
    exit
    selective
        mldp
            no shutdown
        exit
        no auto-discovery-disable
        data-threshold 239.0.0.0/8 1
        data-threshold 239.70.1.0/24 1
        multistream-spmsi 1 create
            group 224.0.0.0/24
            source 192.0.2.0/24
        exit
    exit
    multistream-spmsi 10 create
        group 224.0.0.0/24
        source 192.0.1.0/24
    exit
    exit
    exit
exit

```

3.4.6.11 Policy-based data MDT

A single data MDT can transport one or more IPv4 (C-S, C-G) streams. This allows multiple multicast streams starting from the same VPRN going to the same set of leafs to use a single data MDT. Characteristics of an MDT include:

- a multistream data MDT can carry IPv4 only
- a multistream data MDT can coexist with a single-group data MDT
- a default MDT must be configured

To create a multistream MDT data, an MDA data policy must be configured in the context of MVPN on the source node. This policy maps multiple (C-S, C-G) streams to a single data MDT. Because this configuration is per MVPN, multiple VPNs can have identical policies configured, each for its own VPN context.

When a multicast stream is mapped to a multistream data MDT policy, the data traverses the default MDT first. The data delay timer is used to switch the data from the default MDT to the multistream data MDT.

When the multistream data MDT is deleted, the traffic switches back to the default MDT.

In some cases when a new multistream data MDT is configured which is better suited, some streams may prefer this new multistream data MDT. To switch, the traffic switches to the default MDT and then to the new multistream data MDT.

MDT data can be configured as SSM or ASM.

There can be multiple multistream policies assigned to a single data MDT configuration. In this case, the policy acts as a link list, the first (lowest index) matched multistream policy is used for that specific stream. The following are the rules of matching a multistream data MDT to (C-S, C-G) mapping on a source PE (in order):

1. The multistream policy is evaluated to enable the feature in service when per-(C-S, C-G) configuration is present, starting from the lowest numerical policy index (only entries that are not shutdown are evaluated). The first multistream data MDT (the lowest policy index) the (C-S, C-G) maps to is selected.
2. If (C-S, C-G) does not map to any of the multistream data MDTs, per-(C-S, C-G) single data MDTs are used if a (C-S, C-G) maps to an existing MDT based on data-thresholds.
3. The default MDT is used if no policy matches to a data MDT.
4. When a (C-S, C-G) arrives, it start on the default MDT but can switch to a data MDT when the data delay interval expires. It does not check the data-threshold before switching.
5. When going from one multistream data MDT to a more suitable one, the traffic first switches to the default MDT and then switch to the new multistream data MDT based on the data-delay-interval.

When a data MDT tunnel fails, all (C-S, C-G)s using this multistream data MDT move to the default MDT, and the groups move back to the data MDT when it is restored.

3.4.7 MVPN (NG-MVPN) upstream multicast hop fast failover

MVPN upstream PE or P node fast failover detection method is supported with RSVP P2MP I-PMSI only. A receiver PE achieves fast upstream failover based on the capability to subscribe multicast flow from multiple UMH nodes and the capability to monitor the health of the upstream PE and intermediate P nodes using an unidirectional multipoint BFD session running over the provider tunnel.

A receiver PE subscribes multicast flow from multiple upstream PE nodes to have active redundant multicast flow available during failure of primary flow. Active redundant multicast flow from standby upstream PE allows instant switchover of multicast flow during failure of primary multicast flow.

Faster detection of multicast flow failure is achieved by keeping track of unidirectional multipoint BFD sessions enabled on the provider tunnel. Multi-point BFD sessions must be configured with 10 ms transmit interval on sender (root) PE to achieve sub-50ms fast failover on receiver (leaf) PE.

Configure the **tunnel-status** command option of the following command on the receiver PE for upstream fast failover.

```
configure service vprn mvpn umh-selection
```

Primary and standby upstream PE pairs must be configured on the receiver PE to enable active redundant multicast flow to be received from the standby upstream PE.

3.4.8 Multicast VPN extranet

Multicast VPN extranet distribution allows multicast traffic to flow across different routing instances. A routing instance that received a PIM/IGMP JOIN but cannot reach source of multicast source directly within its own instance is selected as receiver routing instance (receiver C-instance). A routing instance that has source of multicast stream and accepts PIM/IGMP JOIN from other routing instances is selected as source routing instance (source C-instance). A routing instance that does not have either source or receivers but is used in the core is selected as a transit instance (transit P-instance). The following subsections detail supported functionality.

3.4.8.1 Multicast extranet for Rosen MVPN for PIM SSM

Multicast extranet is supported for Rosen MVPN with MDT SAFI. Extranet is supported for IPv4 multicast stream for default and data MDTs (PIM and IGMP).

The following extranet cases are supported:

- local replication into a receiver VRF from a source VRF on a source PE
- transit replication from a source VRF onto a tunnel of a transit core VRF on a source PE

A source VRF can replicate its streams into multiple core VRFs as long as any specific stream from source VRF uses a single core VRF (the first tunnel in any core VRF on which a join for the stream arrives). Streams with overlapping group addresses (same group address, different source address) are supported in the same core VRF.

- remote replication from source or transit VRF into one or more receiver VRFs on receiver PEs
- multiple replications from multiple source or transit VRFs into a receiver VRF on receiver PEs

Rosen MVPN extranet requires routing information exchange between the source VRF and the receiver VRF based on route export or import policies:

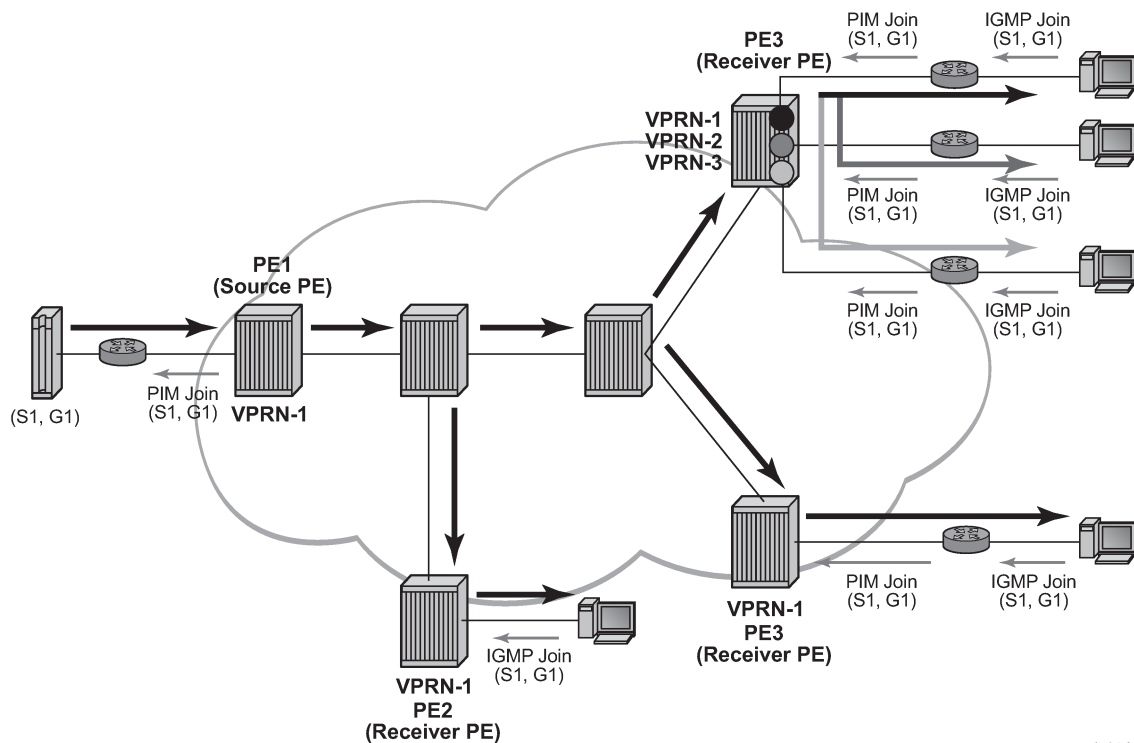
- Routing information for multicast sources must be exported using an RT export policy from the source VRF instance.
- Routing information must be imported into the receiver or transit VRF instance using an RT import policy.

The following are restrictions:

- The source VRF instance and receiver VRF instance of extranet must exist on a common PE node (to allow local extranet mapping).
- SSM translation is required for IGMP (C-*, C-G).
- An I-PMSI route cannot be imported into multiple VPRNs, and NG-MVPN routes do not need to be imported into multiple VPRNs.

In [Figure 31: Multicast VPN traffic flow](#), VPRN-1 is the source VPRN instance and VPRN-2 and VPRN-3 are receiver VPRN instances. The PIM/IGMP JOIN received on VPRN-2 or VPRN-3 is for (S1, G1) multicast flow. Source S1 belongs to VPRN-1. Because of the route export policy in VPRN-1 and the import policy in VPRN-2 and VPRN-3, the receiver host in VPRN-2 or VPRN-3 can subscribe to the stream (S1, G1).

Figure 31: Multicast VPN traffic flow



al_0164

3.4.8.2 Multicast extranet for NG-MVPN for PIM SSM

Multicast extranet is supported for ng-MVPN with IPv4 RSVP-TE and mLDP I-PMSIs and S-PMSIs including (C-*, C-*) S-PMSI support where applicable. Extranet is supported for IPv4 C-multicast traffic (PIM/IGMP joins).

The following extranet cases are supported:

- local replication into a receiver C-instance MVPNs on a source PE from a source P-instance MVPN
- remote replication from P-instance MVPN into one or more receiver C-instance MVPNs on receiver PEs
- multiple replications from multiple source/transit P-instance MVPNs into a receiver C-instance MVPN on receiver PEs



Note: Transit replication on Source PE is not supported.

Multicast extranet for ng-MVPN, similarly to extranet for Rosen MVPN, requires routing information exchange between source ng-MVPN and receiver ng-MVPN based on route export and import policies. Routing information for multicast sources is exported using an RT export policy from a source ng-MVPN instance and imported into a receiver ng-MVPN instance using an RT import policy. S-PMSI/I-PMSI establishment and C-multicast route exchange occurs in a source ng-MVPN P-instance only (import and export policies are not used for MVPN route exchange). Sender-only functionality must not be enabled for the source/transit ng-MVPN on the receiver PE. It is recommended to enable receiver-only functionality on a receiver ng-MVPN instance.

The following are restrictions:

- Source P-instance MVPN and receiver C-instance MVPN must reside on the receiver PE (to allow local extranet mapping).
- SSM translation is required for IGMP (C-*, C-G).

3.4.8.3 Multicast extranet with per-group mapping for PIM SSM

In some deployments, such as IPTV or wholesale multicast services, it may be desirable to create one or more transit MVPNs to optimize delivery of multicast streams in the provider core. The following figures show a deployment model example:

- [Figure 32: Source PE transit replication and receiver PE per-group SSM extranet mapping \(MD-CLI\)](#)
- [Figure 33: Source PE transit replication and receiver PE per-group SSM extranet mapping \(classic CLI\)](#)

Figure 32: Source PE transit replication and receiver PE per-group SSM extranet mapping (MD-CLI)

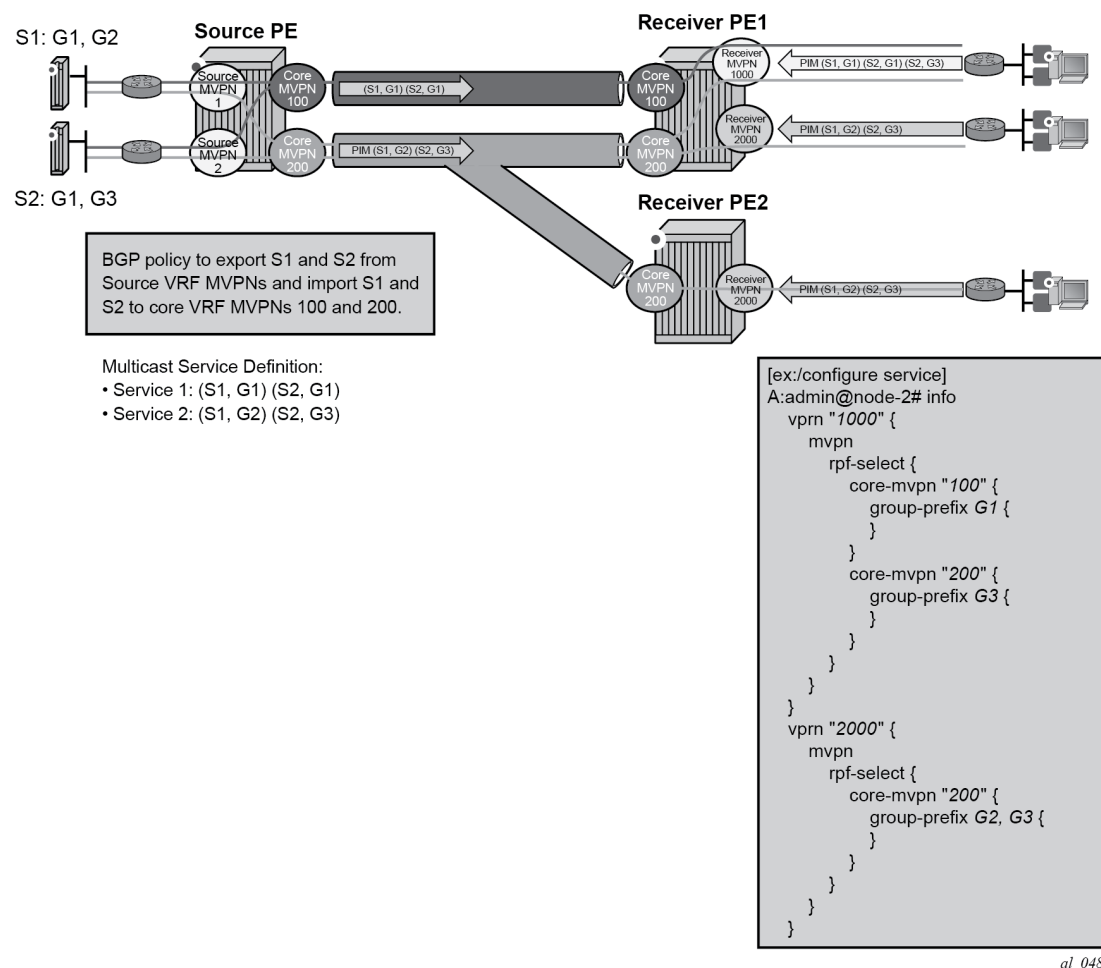
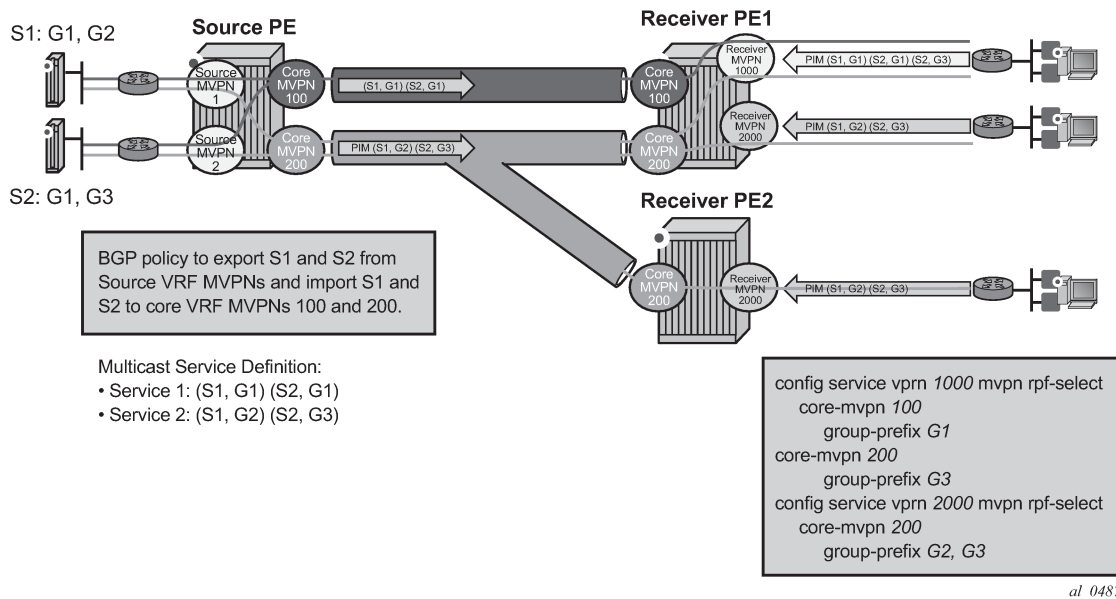


Figure 33: Source PE transit replication and receiver PE per-group SSM extranet mapping (classic CLI)



The architecture displayed in the preceding figures requires a source routing instance MVPN to place its multicast streams into one or more transit core routing instance MVPNs (each stream mapping to a single transit core instance only). It also requires receivers within each receiver routing instance MVPN to know which transit core routing instance MVPN they need to join for each of the multicast streams. To achieve this functionality, transit replication from a source routing instance MVPN onto a tunnel of a transit core routing instance MVPN on a source PE (see earlier sub-sections for MVPN topologies supporting transit replication on source PEs) and per-group mapping of multicast groups from receiver routing instance MVPNs to transit core routing instance MVPNs (as defined below) are required.

For per-group mapping on a receiver PE, the user must configure a receiver routing instance MVPN per-group mapping to one or more source/transit core routing instance MVPNs. The mapping allows propagation of PIM joins received in the receiver routing instance MVPN into the core routing MVPN instance defined by the map. All multicast streams sourced from a single multicast source address are always mapped to a single core routing instance MVPN for a specific receiver routing instance MVPN (multiple receiver MVPNs can use different core MVPNs for groups from the same multicast source address). If per-group map in receiver MVPN maps multicast streams sourced from the same multicast source address to multiple core routing instance MVPNs, then the first PIM join processed for those streams selects the core routing instance MVPN to be used for all multicast streams from a specific source address for this receiver MVPN. PIM joins for streams sourced from the source address not carried by the selected core VRF MVPN instance remains unresolved. When a PIM join or prune is received in a receiver routing instance MVPN with per-group mapping configured, if no mapping is defined for the PIM join's group address, non-extranet processing applies when resolving how to forward the PIM join/prune.

The main attributes for per-group SSM extranet mapping on receiver PE include support for:

- Rosen MVPN with MDT SAFI. RFC6513/6514 NG-MVPN with IPv4 RSVP-TE/mLDP in P-instance (a P-instance tunnel must be in the same VPRN service as multicast source)
- IPv4 PIM SSM
- IGMP (C-S, C-G), and for IGMP (C-*, C-G) using SSM translation
- a receiver routing instance MVPN to map groups to multiple core routing instance MVPNs

- in-service changes of the map to a different transit/source core routing instance (this is service affecting)

The following are restrictions:

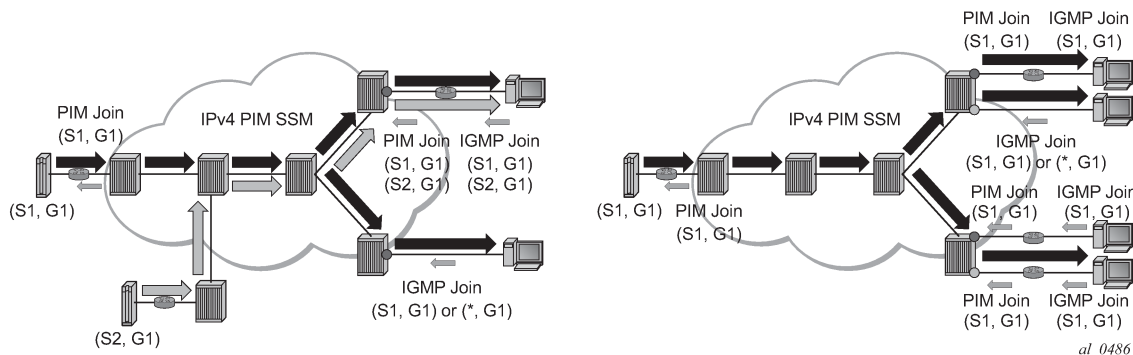
- When a receiver routing instance MVPN is on the same PE as a source routing instance MVPN, basic extranet functionality and not per-group (C-S, C-G) mapping must be configured (extranet from receiver routing instance to core routing instance to source routing instance on a single PE is not supported).
- Local receivers in the core routing instance MVPN are not supported when per-group mapping is deployed.
- Receiver routing instance MVPN that has per-group mapping enabled cannot have tunnels in its OIF lists.
- Per-group mapping is blocked if GRT/VRF extranet is configured.

3.4.8.4 Multicast GRT-source/VRF-receiver extranet with per group mapping for PIM SSM

Multicast GRT-source/VRF-receiver (GRT/VRF) extranet with per-group mapping allows multicast traffic to flow from GRT into VRF MVPN instances. A VRF routing instance that received a PIM/IGMP join but cannot reach the source multicast stream directly within its own instance is selected as receiver routing instance. A GRT instance that has sources of multicast streams and accepts PIM joins from other VRF MVPN instances is selected as source routing instance.

Figure 34: GRT/VRF extranet shows an example deployment.

Figure 34: GRT/VRF extranet



Routing information is exchanged between GRT and VRF receiver MVPN instances of extranet by enabling `grt-extranet` under a receiver MVPN PIM configuration for all or a subset of multicast groups. When enabled, multicast receivers in a receiver routing instances can subscribe to streams from any multicast source node reachable in GRT source instance.

The main feature attributes are:

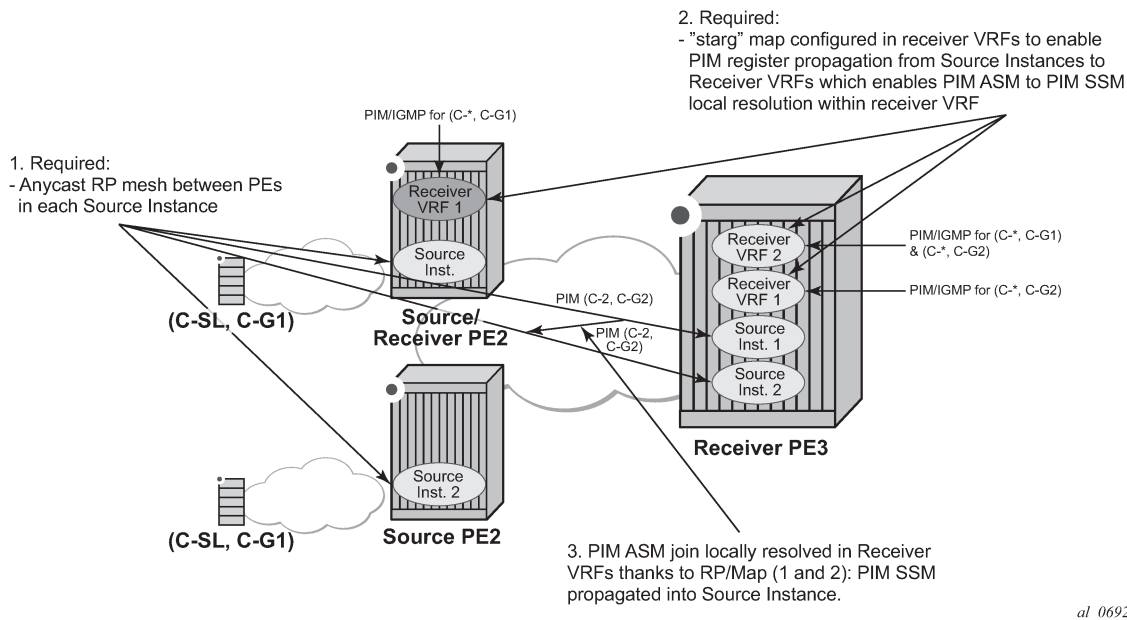
- GRT/VRF extranet can be performed on all streams or on a configured group of prefixes within a receiver routing instance.
- GRT instance requires Classic Rosen multicast.
- IPv4 PIM joins are supported in receiver VRF instances.
- Local receivers using IGMP: (C-S, C-G) and (C-*, C-G) using SSM translation are supported.
- The feature is blocked if a per-group mapping extranet is configured in receiver VRF.

3.4.8.5 Multicast extranet with per-group mapping for PIM ASM

Multicast extranet with per-group mapping for PIM ASM allows multicast traffic to flow from a source routing instance to a receiver routing instance when a PIM ASM join is received in the receiver routing instance.

Figure 35: Multicast extranet with per group PIM ASM mapping depicts PIM ASM extranet map support.

Figure 35: Multicast extranet with per group PIM ASM mapping



PIM ASM extranet is achieved by local mapping from the receiver to source routing instances on a receiver PE. The mapping allows propagation of anycast RP PIM register messages between the source and receiver routing instances over an auto-created internal extranet interface. This PIM register propagation allows the receiver routing instance to resolve PIM ASM joins to multicast sources and to propagate PIM SSM joins over an auto-created extranet interface to the source routing instance. PIM SSM joins are then propagated toward the multicast source within the source routing instance.

The following MVPN topologies are supported:

- Rosen MVPN with MDT SAFI: a local replication on a source PE and multiple-source/multiple-receiver replication on a receiver PE
- RFC 6513/6514 NG-MVPN (including RFC 6625 (C-*, C-*) wildcard S-PMSI): a local replication on a source PE and a multiple source/multiple receiver replication on a receiver PE
- Extranet for GRT-source/VRF receiver with a local replication on a source PE and a multiple-receiver replication on a receiver PE
- Locally attached receivers are supported without SSM mapping.

To achieve the extranet replication:

- Configure in one of the following contexts, as applicable, the local PIM ASM mapping on a receiver PE from a receiver routing instance to a source routing instance.

– **MD-CLI**

```
configure service vprn mvpn rpf-select core-mvpn
configure service vprn pim ipv4 grt-extranet
```

– **classic CLI**

```
configure service vprn mvpn rpf-select core-mvpn
configure service vprn pim grt-extranet
```

- an anycast RP mesh between source and receiver PEs in the source routing instance

The following are restrictions:

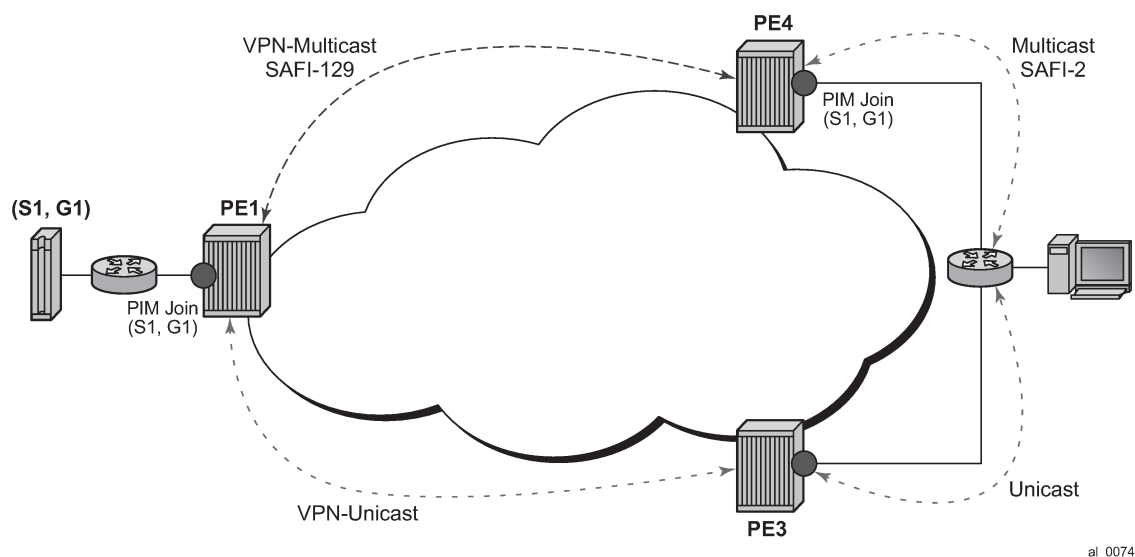
- This feature is supported for IPv4 multicast.
- The multicast source must reside in the source routing instance the ASM map points to on a receiver PE (the deployment of transit replication extranet from source instance to core instance on Source PE with ASM map extranet from receiver instance to core instance on a receiver PE is not supported).
- A specific multicast group can be mapped in a receiver routing instance using either PIM SSM mapping or PIM ASM mapping but not both.
- A specific multicast group cannot map to multiple source routing instances.

3.4.9 Non-congruent unicast and multicast topologies for multicast VPN

The users that prefer to keep unicast and multicast traffic on separate links in a network have the option to maintain two separate instances of the route table (unicast and multicast) per VPRN.

Multicast BGP can be used to advertise separate multicast routes using Multicast NLRI (SAFI 2) on PE-CE link within VPRN instance. Multicast routes maintained per VPRN instance can be propagated between PE-PE using BGP Multicast-VPN NLRI (SAFI 129).

Figure 36: Incongruent multicast and unicast topology for non-overlapping traffic links



SR OS supports option to perform RPF check per VPRN instance using multicast route table, unicast route table or both.

Non-congruent unicast and multicast topology is supported with NG-MVPN. Draft Rosen is not supported.

3.4.10 Automatic discovery of Group-to-RP mappings (auto-RP)

Auto-RP is a proprietary group discovery and mapping mechanism for IPv4 PIM that is described in cisco-ipmulticast/pim-autorp-spec, *Auto-RP: Automatic discovery of Group-to-RP mappings for IP multicast*. The functionality is similar to the IETF standard bootstrap router (BSR) mechanism that is described in RFC 5059, *Bootstrap Router (BSR) Mechanism for Protocol Independent Multicast (PIM)*, to dynamically learn about the availability of Rendezvous Points (RPs) in a network. Use the following command to configure the router as an RP-mapping agent that listens to the CISCO-RP-ANNOUNCE (224.0.1.39) group and caches the announced mappings:

- **MD-CLI**

```
configure router pim rp ipv4 auto-rp-discovery
```

- **classic CLI**

```
configure router pim rp auto-rp-discovery
```

The RP-mapping agent then periodically sends out RP-mapping packets to the CISCO-RP-DISCOVERY (224.0.1.40) group. SR OS supports version 1 of the auto-RP specification, so the ability to deny RP-mappings by advertising negative group prefixes is not supported.

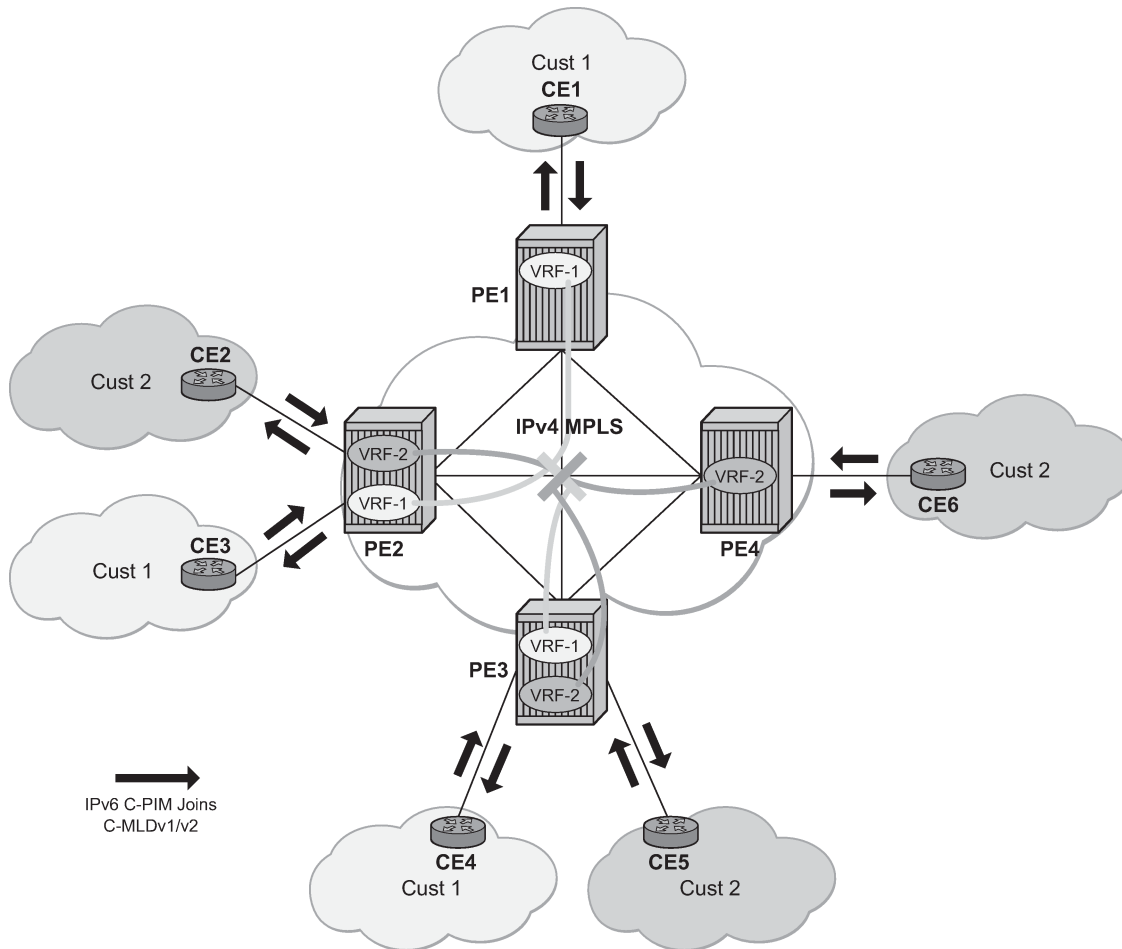
PIM dense-mode (PIM-DM) as described in RFC 3973, *Protocol Independent Multicast - Dense Mode (PIM-DM): Protocol Specification (Revised)*, is used for the auto-RP groups to support multihoming and redundancy. The RP-mapping agent supports announcing, mapping, and discovery functions; candidate RP functionality is not supported.

Auto-RP is supported for IPv4 in multicast VPNs and in the global routing instance. Either BSR or auto-RP for IPv4 can be configured; the two mechanisms cannot be enabled together. BSR for IPv6 and auto-RP for IPv4 can be enabled together. In a multicast VPN, auto-RP cannot be enabled together with sender-only or receiver-only multicast distribution trees (MDTs), or wildcard S-PMSI configurations that could block flooding.

3.4.11 IPv6 MVPN support

IPv6 multicast support in SR OS allows the users to offer customers IPv6 multicast MVPN service. IPv4 mLDP or RSVP-TE core is used to carry IPv6 c-multicast traffic inside IPv4 mLDP or RSVP-TE provider tunnels (p-tunnels). The IPv6 customer multicast on a specific MVPN can be blocked, enabled on its own or in addition to IPv4 multicast per PE or per interface. When both IPv4 and IPv6 multicast is enabled for a specific MVPN, a single tree is used to carry both IPv6 and IPv4 traffic. [Figure 37: IPv6 MVPN example](#) shows an example of a user with IPv4 MPLS backbone providing IPv6 MVPN service to Customer 1 and Customer 2.

Figure 37: IPv6 MVPN example



al_0168

SR OS IPv6 MVPN multicast implementation provides the following functionality:

- IPv6 C-PIM-SM (ASM and SSM)
- MLDv1 and MLDv2
- SSM mapping for MLDv1
- I-PMSI and S-PMSI using IPv4 P2MP mLDP p-tunnels
- I-PMSI and S-PMSI using IPv4 P2MP RSVP p-tunnels
- BGP auto-discovery
- PE-PE transmission of C-multicast routing using BGP mvpn-ipv6 address family
- IPv6 BSR/RP functions on functional par with IPv4 (auto-RP using IPv4 only)
- Embedded RP
- Inter-AS Option A

The following known restrictions exist for IPv6 MVPN support:

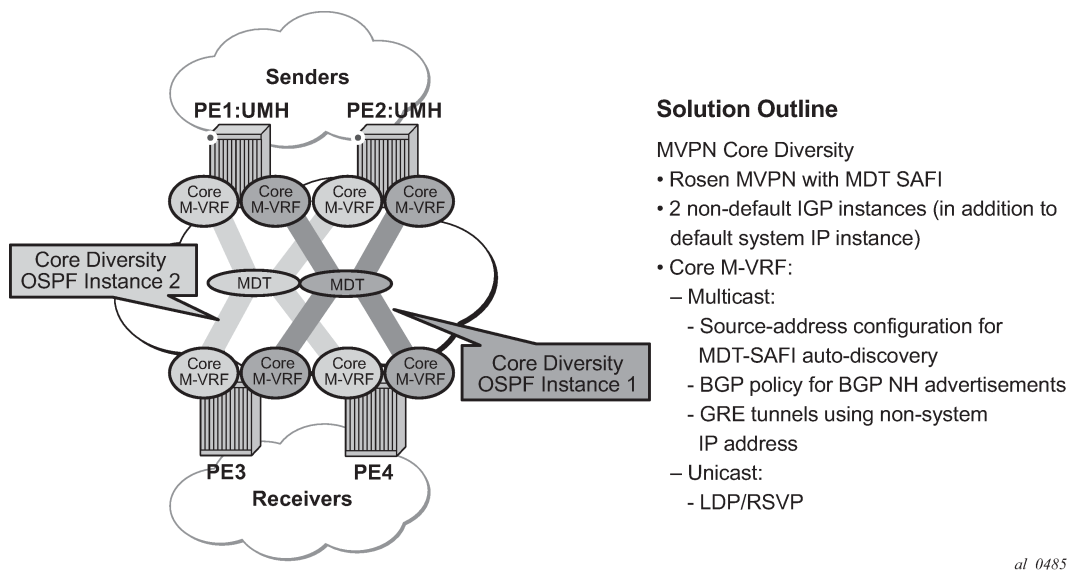
- Non-congruent topologies are not supported.

- IPv6 is not supported in MCAC.
- If IPv4 and IPv6 multicast is enabled, per-MVPN multicast limits apply to entire IPv4 and IPv6 multicast traffic as it is carried in a single PMSI. For example IPv4 AND IPv6 S-PMSIs are counted against a single S-PMSI maximum per MVPN.
- IPv6 Auto-RP is not supported.

3.4.12 Multicast core diversity for Rosen MDT_SAFI MVPNs

Figure 38: Multicast core diversity depicts Rosen MVPN core diversity deployment:

Figure 38: Multicast core diversity



al_0485

Core diversity allows the user to optionally deploy multicast MVPN in either default IGP instance or one of two non-default IGP instances to provide, for example, topology isolation or different level of services. The following describes main feature attributes:

- Rosen MVPN IPv4 multicast with MDT SAFI is supported with default and data MDTs.
- Rosen MVPN can use a non-default OSPF or ISIS instance (using their loopback addresses instead of a system address).
- Up to 3 distinct core instances are supported: system + 2 non-default OSPF instances – referred as “red” and “blue” below.
- The BGP Connector also uses non-default OSPF loopback as NH, allowing Inter-AS Option B/C functionality to work with Core diversity as well.
- The feature is supported with CSC-VPRN.

On source PEs (PE1: UMH, PE2: UMH in Figure 38: Multicast core diversity), an MVPN is assigned to a non-default IGP core instance as follows:

1. MVPN is statically pointed to use one of the non-default “red”/“blue” IGP instances loopback addresses as source address instead of system loopback IP.

2. MVPN export policy is used to change unicast route next-hop VPN address (no longer required as of SR OS Release 12.0.R4 - BGP Connector support for non-default instances).

The above configuration ensures that MDT SAFI and IP-VPN routes for the non-default core instance use non-default IGP loopback instead of system IP. This ensures PIM advertisement/joins run in the correct core instance and GRE tunnels for multicast can be set-up using and terminated on non-system IP.

If BGP export policy is used to change unicast route next-hop VPN address, unicast traffic must be forwarded in non-default "red" or "blue" core instance LDP or RSVP (terminating on non-system IP) must be used. GRE unicast traffic termination on non-system IP is not supported, and any GRE traffic arriving at the PE in "blue", "red" instances destined for non-default IGP loopback IP is forwarded to CPM (ACL or CPM filters can be used to prevent the traffic from reaching the CPM). This limitation does not apply if BGP connector attribute is used to resolve the multicast route.

No configuration is required on non-source PEs.

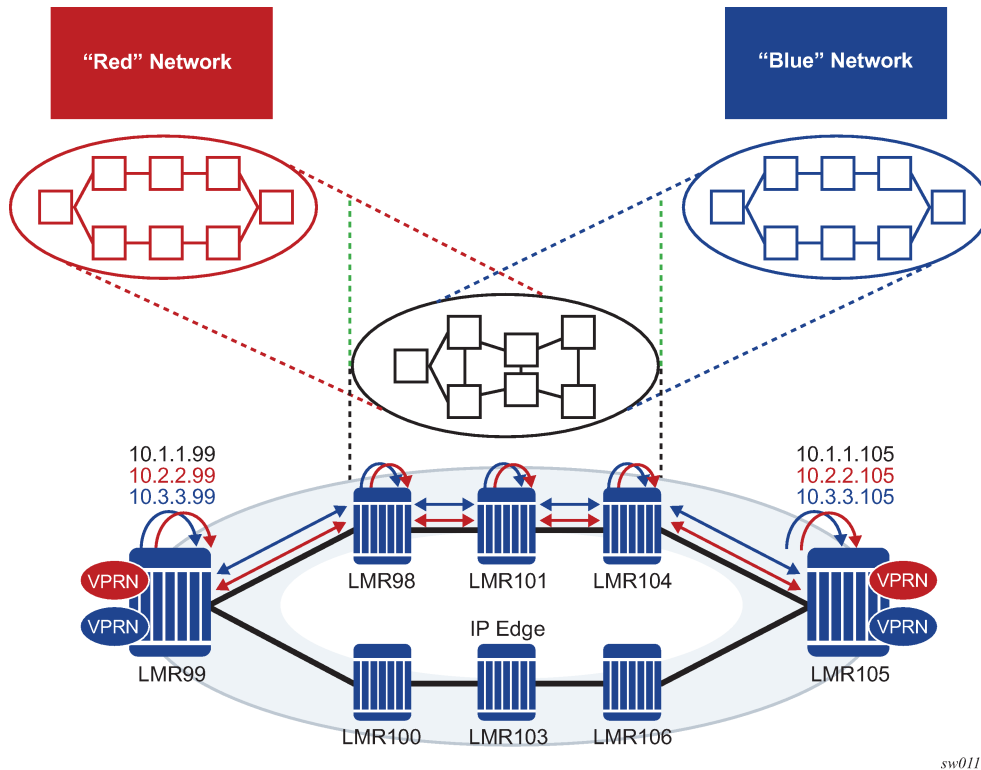
The following are feature restrictions:

- VPRN instance must be shutdown to change the mdt-safi source-address. The CLI rollback that includes change of the auto-discovery is therefore service impacting.
- To reset mdt-safi source-address to system IP, the user must first execute no auto-discovery (or auto-discovery default) then auto-discovery mdt-safi
- Configuring system IP as a source-address consumes one of the 2 IP addresses allowed, therefore it should not be done.
- The users must configure correct IGP instance loopback IP addresses within Rosen MVPN context and must configure correct BGP policies (Before Release 12.0.R4) for the feature to operate as expected. There is no verification that the address entered for MVPN provider tunnel source-address is such an address or is not a system IP address.

3.4.13 NG-MVPN core diversity

The following figure shows an operational example of logical networks using multi-instance IGP.

Figure 39: Logical networks using multi-instance IGP



The SR OS is used in multi-instance IGP as a virtualization or migration strategy in numerous cases. One application is as a virtual LSR core, where various topologies are created using separate IGP instances. Specifically, the migration to SR solutions requires the deployment of multi-instance IGP with service migrations. The objective is to more cleanly segregate protocols and service bindings to specific routing instances. NG-MVPN does not currently allow the use of non-system loopbacks for PMSI (for example, MVPN address family).

The ability to support binding MVPN with mLDP tunnels to different loopback interfaces is one of its main uses. In addition, assigning these loopback interfaces to different IGP instances creates parallel NG-MVPN services, each running over a separate IGP instance.

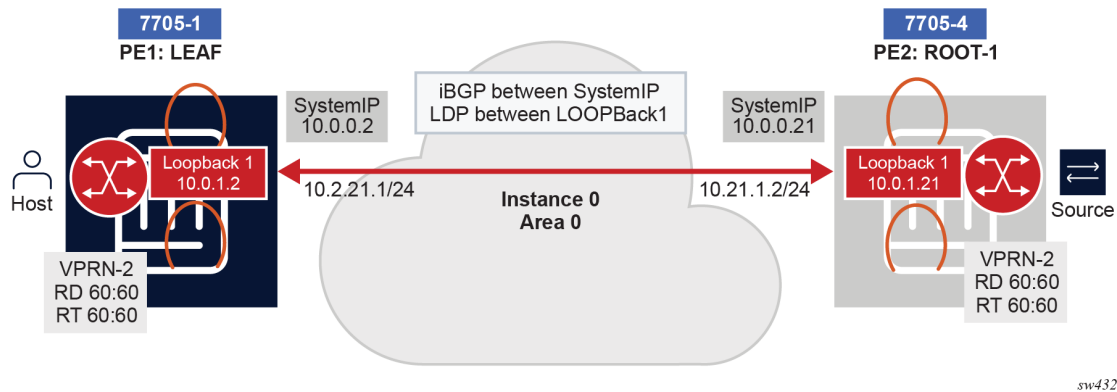
NG-MVPN core diversity has the following main components, which, when used in combination, create parallel NG-MVPN services on different IGP instances:

- the ability to advertise an MVPN route via a loopback interface and generate an mLDP FEC with the loopback interface
- the ability to create multiple IGP instances and assign a loopback to each instance

3.4.13.1 NG-MVPN to loopback interface

The following figure shows an NG-MVPN setup via a loopback interface.

Figure 40: NG-MVPN setup via loopback interface



The system IP is typically used for management purposes, so it may be more appropriate to create services to a separate loopback interface. Because of security concerns, many operators avoid creating services to the system IP address.

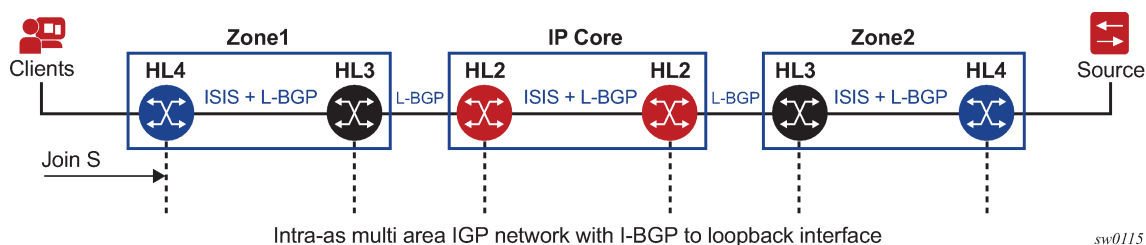
MVPN routes can be advertised with a next-hop specific loopback, which can be achieved using the following methods:

- configure the iBGP session to the loopback interface of the peer and use the corresponding local loopback for the local address
- configure the iBGP session to the system IP but use policies to change the next hop of the AD route to the corresponding loopback interface

After the AD routes are advertised via the loopback interface as the next hop, PIM generates an mLDP FEC for the loopback interface.

The preceding configuration allows an NG-MVPN to be established via a specific loopback.

Figure 41: Intra-AS basic opaque FEC to loopback interface



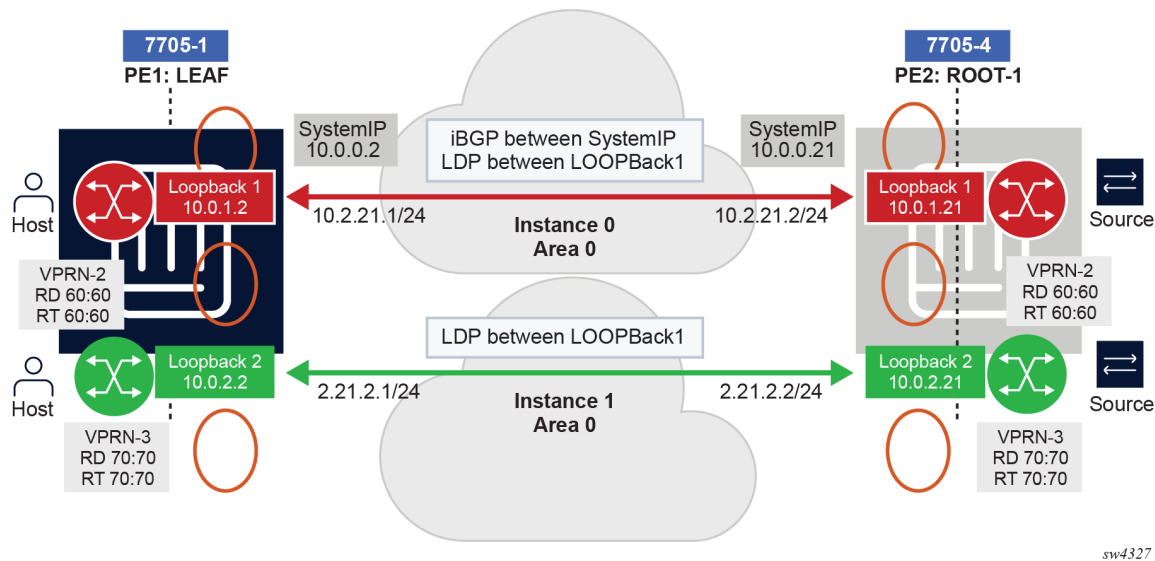
It should be noted that this setup works in a single area or multi-area through an ABR.

3.4.13.2 NG-MVPN core diversity

In core diversity, MVPN services on the LER are bound to a domain by advertising MP-BGP routes, with the next hop set to the loopback address; the receiving LERs resolve the BGP next hops based on the corresponding IGP instance. All subsequent MVPN BGP routing exchanges must set and resolve the BGP next hop with the configured non-system loopback.

The following figure shows that red and blue networks correspond to separate IGP instances tied to separate loopback interfaces.

Figure 42: Core diversity with parallel NG-MVPN services on parallel IGP instances



With this feature, an MP-BGP next hop resolves to a link LDP label that is indirectly associated with an IGP instance. Services that advertise BGP routes with a next-hop red loopback result in traffic flowing over the red network using LDP, and the blue loopback results in traffic flowing over the blue network.

Most importantly, the common routes in each IGP instance must have a unique and active label. This is required to ensure that the same route advertised in two different IGP domains does not resolve to the same label. The LDP *local-lsr-id* can be used to ensure FECs and label mapping are advertised using the correct instance of IGP.

Core diversity allows an operator to optionally deploy NG-MVPN in either the default IGP instance or one of the non-default IGP instances, which provides, for example, topology isolation or different level of services. The following are the main feature attributes:

- NG-MVPN can use IPv4 or IPv6 multicast.
- NG-MVPN can use a non-default OSPF or IS-IS instance.



Note: Core diversity is accomplished by using the loopback addresses instead of a system address.

- The BGP connector also uses non-default OSPF loopback as the next hop, using the following methods:
 - setting the BGP local-address to a loopback interface IP address
 - creating a routing policy to set the NH of an MVPN AD route to the corresponding loopback IP
- RSVP-TE/LDP transport tunnels also use non-systemIP loopback for session creation. As an example, RSVP-TE p2mp LSPs would be created to non systemIP loopbacks and mLDP creates a session via the local LSR ID.
- The *source-address* for the **mvpn auto-discover** command default must be included.

On the source PEs, an NG-MVPN is assigned to a non-default IGP core instance, as follows:

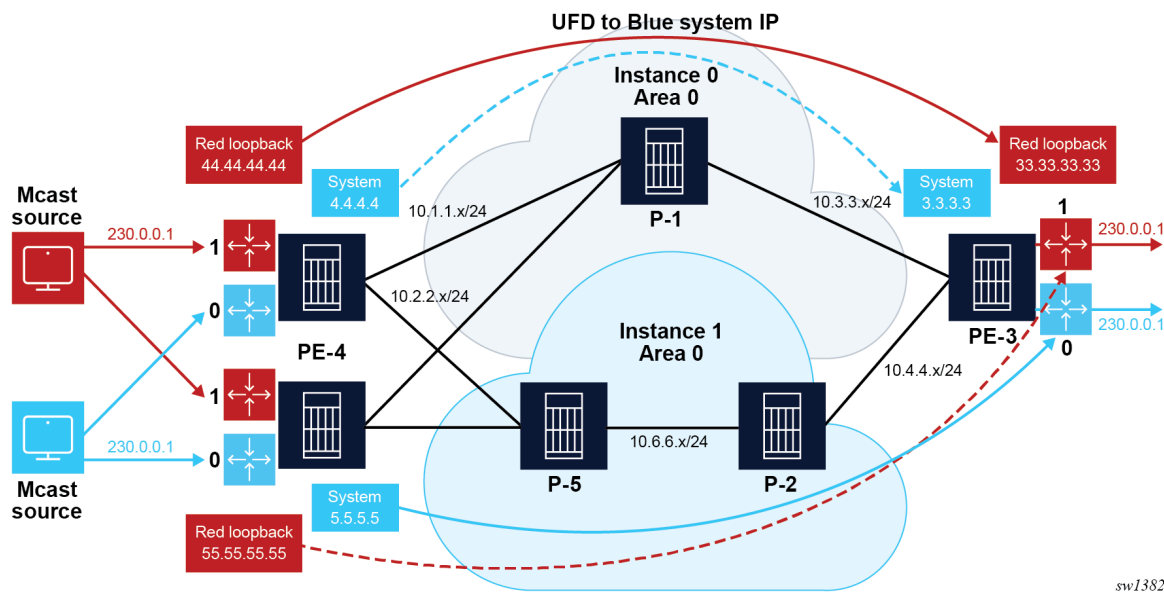
- NG-MVPN is statically pointed to use one of the non-default red or blue IGP instance loopback addresses as the source address instead of the system loopback IP.
- The MVPN export policy is used to change the unicast route next-hop VPN address (BGP connector support for non-default instances).
- Alternatively, the BGP local address can be set to the correct loopback interface assigned for the specific instance.

The preceding configuration ensures that MVPN-IPv4/v6 and IP-VPN routes for the non-default core instance use non-default IGP loopback instead of system IP. This ensures MVPN-IPv4/v6 advertisement/joins run in the correct core instance and mLDP and P2MP RSVP tunnels (I-PMSI and S-PMSI) for multicast can be set-up using and terminating on non-system IP.

If BGP export policy is used to change unicast route next-hop VPN address, then unicast traffic must be forwarded in non-default Red or Blue core instance LDP or RSVP (terminating on non-system IP) must be used.

3.4.13.3 P2MP RSVP-TE core diversity with UFD for UMH redundancy

Figure 43: P2MP RSVP-TE core diversity with UFD for UMH redundancy



Each NG-MVPN can be established over a separate IGP instance for core diversity support. In [Figure 43: P2MP RSVP-TE core diversity with UFD for UMH redundancy](#), there are two IGP instances, Blue and Red, each having its own dedicated NG-MVPN service. The Blue plane is identified with the system IP address, and Red plane is identified with a loopback IP address. MVPN Autodiscovery (AD) routes are advertised and installed accordingly based on their NLRI in each instance. To advertise these MVPN AD routes in the corresponding planes, route policies are used to change the next hops of the AD routes. For example, for the Blue plane the AD routes are advertised by default with the system IP address as their next hop, so no route policy is necessary. In the Red plane, a route policy can change the next hop of the corresponding NG-MVPN AD routes to use the loopback IP address.

Core diversity also supports UMH redundancy solutions. In the case of P2MP RSVP-TE, which uses UFD for source UMH failure, the UFD packets are generated for each IGP instance (Blue, Red) accordingly, with the correct source and destination IP address that is part of that IGP instance.

3.4.13.4 UFD packet generation

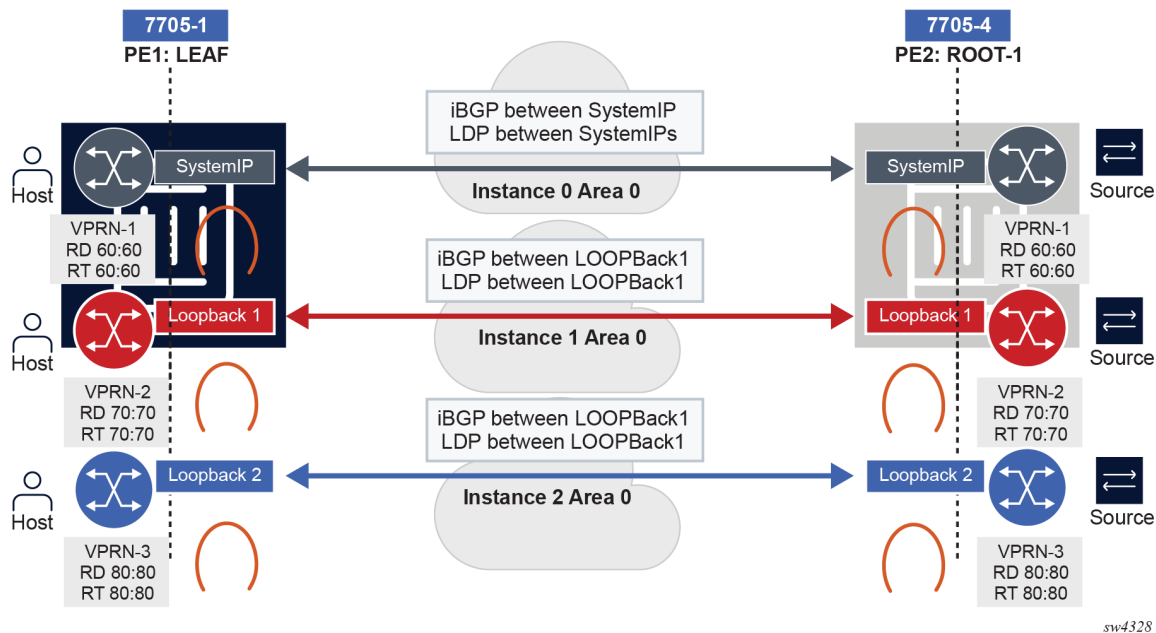
SR OS generates the UFD packets with the correct source and destination IP addresses for each IGP instance used in the core diversity configuration. SR OS uses the NLRI next-hop information of the AD route for the destination IP address and derives the source IP address from the MVPN autodiscovery default source-address loopback for the UFD packets. This ensures the UFD packets are traversing the appropriate core corresponding to their UMH reachability.

For example, in [Figure 43: P2MP RSVP-TE core diversity with UFD for UMH redundancy](#), the Blue plane UFD packets are generated from the root PE to the leaf PE with the source IP address as the local system IP on the root and the destination IP address as the system IP of the leaf. The Red plane generates the UFD packet with the loopback IP address to which it is corresponding to as its source and destination IP address.

3.4.13.5 Configuration example

About this task

Figure 44: Core diversity with parallel NG-MVPN services on parallel IGP instances



In [Figure 44: Core diversity with parallel NG-MVPN services on parallel IGP instances](#), there are three IGP instances, the default Instance 0, Instance 1, and Instance 2. Each instance binds to its own loopback interface. For Instance 0, it is the *SystemIP* system interface used as loopback. LDP, RSVP and MP-BGP need to run between the corresponding loopback associated with each instance.

For example, for the Blue Instance 2, both MP-BGP and LDP need to be configured to its corresponding loopback *loopback2* and the next-hop for BGP MVPN-IP4/IPv6 and the VPN-IP4/IPv6 need to be *loopback2*.

From a configuration point of view, the following steps need to be performed:

Procedure

Step 1. For MLDP, configure LDP with **local-lsr-id** with loopback interface of instance 2 *loopback2* as shown in the following example.

Example

MD-CLI

```
[ex:/configure router "Base" ldp]
A:admin@node-2# info
    interface-parameters {
        interface "loopback2" {
            ipv4 {
            }
        }
        interface "to-root" {
            ipv4 {
                local-lsr-id {
                    interface-name "loopback2"
                }
            }
        }
    }
}
```

Example

classic CLI

```
A:node-2>config>router>ldp# info
...
    interface-parameters
        interface "to-root" dual-stack
            ipv4
                local-lsr-id interface-name "loopback2"
                no shutdown
            exit
            no shutdown
        exit
        interface "loopback2" dual-stack
            ipv4
                no shutdown
            exit
            no shutdown
        exit
    exit
    exit
    targeted-session
    exit
    no shutdown
```

Step 2. Configure the source address for default autodiscovery.

Example

MD-CLI

```
[ex:/configure service vprn "2"]
A:admin@node-2# info
```

```

...
    mvpn {
        auto-discovery {
            type bgp
            source-address LOOPBack1
        }
        vrf-export {
            policy ["vprnexp100"]
        }
    }
[ex:/configure service vprn "3"]
A:admin@node-2# info
...
    mvpn {
        auto-discovery {
            type bgp
            source-address LOOPBack2
        }
        vrf-export {
            policy ["vprnexp101"]
        }
    }
}

```

Example classic CLI

```

A:node-2>config>service#info
vprn 2 name "2" customer 1 create
...
    mvpn
        auto-discovery default source-address LOOPBack1
        vrf-export "vprnexp100"
    exit
exit
vprn 3 name "3" customer 1 create
...
    mvpn
        auto-discovery default source-address LOOPBack2
        vrf-export "vprnexp101"
    exit

```

- Step 3.** Define the community *vprnXXXX* for each VPRN using non-default core-instance and define a policy to tag each VPRN with either a *blue* or *red* standard community attribute.

Example MD-CLI

```

[ex:/configure policy-options]
A:admin@node-2# info
    policy-options {
        community "blue" {
        }
        community "red" {
        }
        community "vprn2" {
            member "target:70:70" { }
        }
        community "vprn3" {
            member "target:80:80" { }
        }
        policy-statement "vprnexp2" {
            entry 10 {

```

```

        from {
            protocol {
                name [direct]
            }
        }
        action {
            action-type accept
            community {
                add ["vprn2" "red"]
            }
        }
    }
}
policy-statement "vprnexp3" {
    entry 10 {
        from {
            protocol {
                name [direct]
            }
        }
        action {
            action-type accept
            community {
                add ["vprn3" "blue"]
            }
        }
    }
}
}

```

Example classic CLI

```

A:node-2>config>router>policy-options# info
-----
community "red"
exit
community "blue"
exit
community "vprn2"
members "target:70:70"
exit
community "vprn3"
members "target:80:80"
exit
policy-statement "vprnexp2"
entry 10
from
protocol direct
exit
action accept
community add "vprn2" "red"
exit
exit
exit
policy-statement "vprnexp3"
entry 10
from
protocol direct
exit
action accept
community add "vprn3" "blue"
exit

```

```

        exit
    exit
-----

```

Step 4. Define a single global BGP policy to change the next hop for *red* and *blue* MVPNs.

Example

MD-CLI

```

[ex:/configure policy-options]
A:admin@node-2# info
  policy-statement "MVPN_CoreDiversity_Exp" {
    entry 10 {
      from {
        community {
          name "red"
        }
      }
      to {
        protocol {
          name [bgp-vpn]
        }
      }
      action {
        action-type accept
        next-hop "@loopback1@"
      }
    }
    entry 20 {
      from {
        community {
          name "blue"
        }
      }
      to {
        protocol {
          name [bgp-vpn]
        }
      }
      action {
        action-type accept
        next-hop "@loopback2@"
      }
    }
  }
}

```

Example

classic CLI

```

A:node-2config>router>policy-options# info
  policy-statement "MVPN_CoreDiversity_Exp"
    entry 10
      from
        community "red"
      exit
      to
        protocol bgp-vpn
      exit
      action accept
        next-hop @loopback1@
      exit
    exit
  exit

```

```

        entry 20
        from
            community "blue"
        exit
        to
            protocol bgp-vpn
        exit
        action accept
            next-hop @loopback2@
        exit
    exit
exit

```

Step 5. Configure the BGP default MVPN export in the group as required.

Example

MD-CLI

```

[ex:/configure router "Base"]
A:admin@node-2# info
    bgp {
        group "mvpn" {
            export {
                policy ["MVPN_CoreDiversity_Exp"]
            }
        }
    }
    mpls {
    }
    rsvp {
    }

```

Example

classic CLI

```

A:node-2>config>router>bgp$ info
-----
        group "mvpn"
            export "MVPN_CoreDiversity_Exp"
        exit
        no shutdown
-----

```

Step 6. Configure each VPRN to use the correct IGP source address and correct VRF export policy.

Example

MD-CLI

```

[ex:/configure service vprn "2" mvpn]
A:admin@node-2# auto-discovery source-address loopback1

*[ex:/configure service vprn "2" mvpn]
A:admin@node-2# vrf-export policy "vprnexp2"

[ex:/configure service vprn "3" mvpn]
A:admin@node-2# auto-discovery source-address loopback2

*[ex:/configure service vprn "3" mvpn]
A:admin@node-2# vrf-export policy "vprnexp3"

```

Example classic CLI

```
*A:node-2>service>vprn 2
  mvpn auto-discovery default source-address loopback1
  vrf-export "vprnexp2"

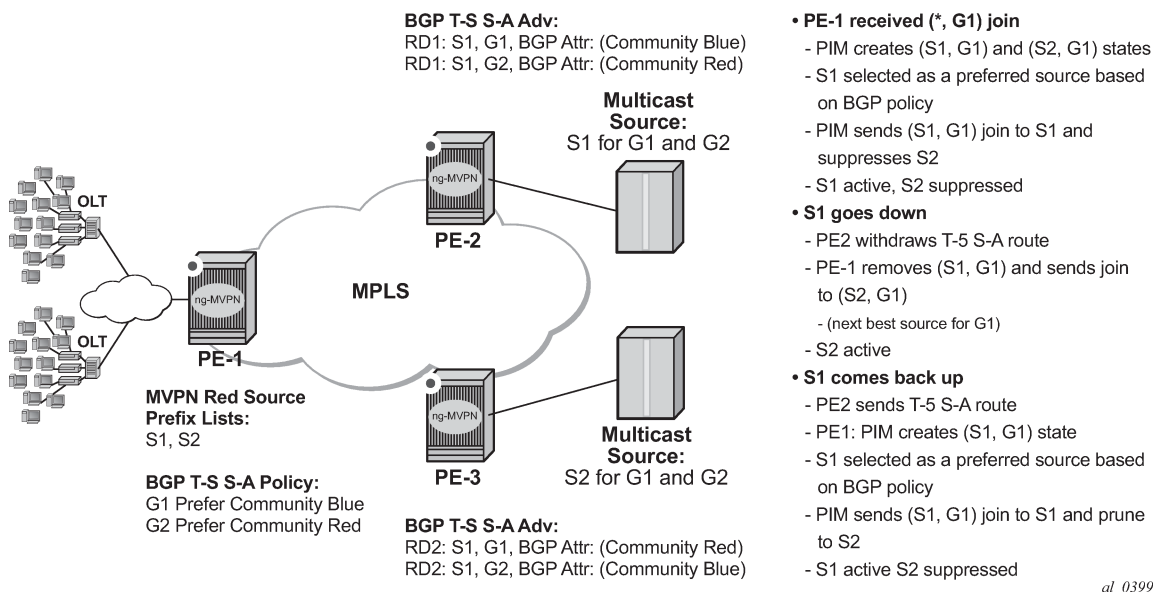
*A:node-2>config>service>vprn 3
  mvpn auto-discovery default source-address loopback2
  vrf-export "vprnexp3"
```

3.4.14 NG-MVPN multicast source geo-redundancy

Multicast source geo-redundancy is targeted primarily for MVPN deployments for multicast delivery services like IPTV. The solution allows the users to configure a list of geographically dispersed redundant multicast sources (with different source IPs) and then, using configured BGP policies, ensure that each Receiver PE (a PE with receivers in its C-instance) selects only a single, most-preferred multicast source for a specific group from the list. Although the data may still be replicated in P-instance (each multicast source sends (C-S, C-G) traffic onto its I-PMSI tree or S-PMSI tree), each Receiver PE only forwards data to its receivers from the preferred multicast source. This allows the users to support multicast source geo-redundancy without the replication of traffic for each (C-S, C-G) in the C-instance while allowing fast recovery of service when an active multicast source fails.

Figure 45: Preferred source selection for multicast source geo-redundancy shows an operational example of multicast source geo-redundancy.

Figure 45: Preferred source selection for multicast source geo-redundancy



The users can configure a list of prefixes for multicast source redundancy per MVPN on Receiver PEs. Up to 8 multicast source prefixes per VPRN are supported. Any multicast source that is not part of the source prefix list is treated as a unique source and automatically joined in addition to joining the most preferred source from the redundant multicast source list.

A Receiver PE selects a single, most-preferred multicast source from the list of pre-configured sources for a specific MVPN during (C-*, C-G) processing as follows:

- A single join for the group is sent to the most preferred multicast source from the user-configured multicast source list. Joins to other multicast sources for a specific group are suppressed. The user can see active and suppressed joins on a Receiver PE. Although a join is sent to a single multicast source only, (C-S, C-G) state is created for every source advertising Type-5 S-A route on the Receiver PE.
- The most preferred multicast source is a reachable source with the highest local preference for Type-5 SA route based on the BGP policy, as described later in this section.
- On a failure of the preferred multicast source or when a new multicast source with a better local preference is discovered, Receiver PE joins the new most-preferred multicast source. The outage experienced depends on how quickly Receiver PE receives Type-5 S-A route withdrawal or loses unicast route to multicast source, and how quickly the network can process joins to the newly selected preferred multicast sources.
- Local multicast sources on a Receiver PE are not subject to the most-preferred source selection, regardless of whether they are part of redundant source list or not.

BGP policy on Type-5 SA advertisements is used to determine the most preferred multicast source based on the best local preference as follows:

- Each Source PE (a PE with multicast sources in its C-instance) tags Type-5 SA routes with a unique standard community attribute using global BGP policy or MVPN vrf-export policy. Depending on multicast topology, the policy may require source-aware tagging in the policy. Either all MVPN routes or Type-5 SA routes only can be tagged in the policy (attribute **mvpn-type 5**). Use the following command to configure the MVPN type to **5**:

– **MD-CLI**

```
configure policy-options policy-statement entry from mvpn-type
```

– **classic CLI**

```
configure router policy-options policy-statement entry from mvpn-type
```

- Each receiver PE has a BGP VRF import policy that sets local preference using match on Type-5 SA routes (attribute **mvpn-type 5**) and standard community attribute value (as tagged by the Source PEs). Using policy statements that also include group address match allows the receiver PEs to select the best multicast source per group. Use the following command to apply the BGP VRF import policy.

```
configure service vprn mvpn vrf-import
```

Use the following command to configure the default action to **accept**; otherwise, all MVPN routes other than those matched by specified entries are rejected:

– **MD-CLI**

```
configure policy-options policy-statement default-action action-type
```

– **classic CLI**

```
configure router policy-options policy-statement default-action
```

In addition, the VRF route target must be configured as a community match condition because the MVPN route target configuration is ignored when the VRF import policy is defined.

Use the following command to configure the VRF route target.

```
configure service vprn mvpn vrf-target
```

Use the following command to configure the VRF import policy.

```
configure service vprn mvpn vrf-import
```

The users can change redundant source list or BGP policy affecting source selection in service. If such a change of the list/policy results in a new preferred multicast source election, make-before-break is used to join the new source and prune the previously best source.

For the correct operations, MVPN multicast source geo-redundancy requires the router:

- To maintain the list of eligible multicast sources on Receiver PEs, Source PE routers must generate Type-5 S-A route even if the Source PE sees no active joins from any receiver for a specific group.
- To trigger a switch from a currently active multicast source on a Receiver PE, Source PE routers must withdraw Type-5 S-A route when the multicast source fails or alternatively unicast route to multicast source must be withdrawn or go down on a Receiver PE.

MVPN multicast source redundancy solutions is supported for the following configurations only. Enabling the feature in unsupported configuration must be avoided.

- NG-MVPN with RSVP-TE or mLDP or PIM with BGP c-multicast signaling in P-instance. Both I-PMSI and S-PMSI trees are supported.
- IPv4 and IPv6 (C-*, C-G) PIM ASM joins in the C-instance.
- Configuring the use of inter-site shared C-trees is optional. Use the following command to configure using inter-site shared C-trees.

```
configure service vprn mvpn intersite-shared
```

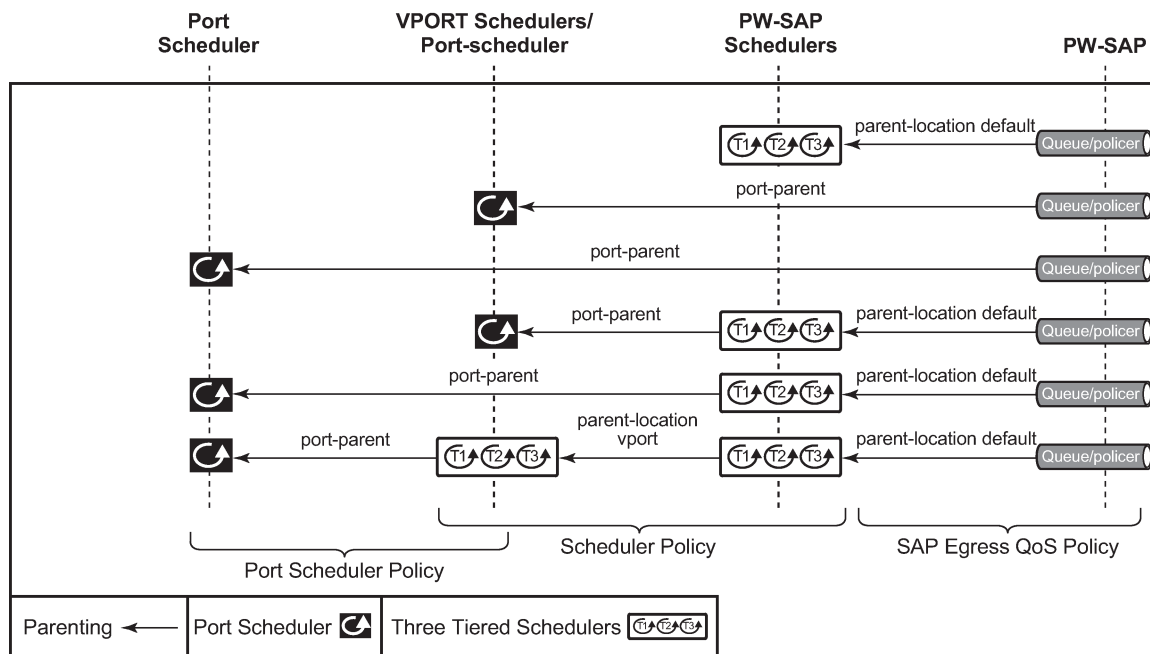
If **intersite-shared** is configured, the user must enable generation of Type-5 S-A routes even in the absence of receivers seen on Source PEs (intersite-shared persistent-type5-adv must be enabled).

- The Source PEs must be configured as a sender-receiver, the Receiver PEs can be configured as a sender-receiver or a receiver-only.
- The RPs must be on the Source PEs side. Static RP, anycast-RP, embedded-RP types are supported.
- UMH redundancy can be deployed to protect Source PE to any multicast source. When deployed, UMH selection is executed independently of source selection after the most preferred multicast source had been chosen. Supported **umh-selection** command options include: **highest-ip**, **hash-based**, **tunnel-status** (not supported for IPv6), and **unicast-rt-pref**.

3.4.15 Multicast core diversity for Rosen MDT SAFI MVPNs

Figure 46: Multicast core diversity shows a Rosen MVPN core diversity deployment.

Figure 46: Multicast core diversity



al_0376

Core diversity allows the users to optionally deploy multicast MVPN in either default IGP instance. or one of two non-default IGP instances to provide; for example, topology isolation or different level of services. The following describes the main feature attributes:

- Rosen MVPN IPv4 multicast with MDT SAFI is supported with default and data MDTs.
- Rosen MVPN can use a non-default OSPF or ISIS instance (using their loopback addresses instead of a system address).
- Up to 3 distinct core instances are supported: system + 2 non-default OSPF instances shown in [Figure 46: Multicast core diversity](#).
- The BGP Connector also uses non-default OSPF loopback as NH, allowing Inter-AS Option B/C functionality to work with Core diversity as well.
- The feature is supported with CSC-VPRN.

On source PEs (PE1: UMH, PE2: UMH in the above picture), an MVPN is assigned to a non-default IGP core instance as follows:

- MVPN is statically pointed to use one of the non-default IGP instances loopback addresses as source address instead of system loopback IP.
- MVPN export policy is used to change unicast route next-hop VPN address.
- BGP Connector support for non-default instances.

The configuration shown above ensures that MDT SAFI and IP-VPN routes for the non-default core instance use non-default IGP loopback instead of system IP. This ensures PIM advertisement/joins run in the correct core instance and GRE tunnels for multicast can be set-up using and terminated on non-system IP. If BGP export policy is used to change unicast route next-hop VPN address instead of BGP Connector attribute-based processing and unicast traffic must be forwarded in non-default core instances 1 or 2, LDP or RSVP (terminating on non-system IP) must be used. GRE unicast traffic termination on non-system IP

is not supported and any GRE traffic arriving at the PE in instances 1 or 2, destined for non-default IGP loopback IP is forwarded to CPM (ACL or CPM filters can be used to prevent the traffic from reaching the CPM).

No configuration is required on the non-source PEs.

Known feature restrictions include:

- VPRN instance must be shut down to change the MDT SAFI source address. The CLI rollback that includes changing the autodiscovery is therefore service impacting.
- To reset the MDT SAFI source address to the system IP, the user must configure no autodiscovery (or autodiscovery default), then autodiscovery MDT-SAFI.
- Configuring the system IP as a source address consumes one of the 2 IP addresses allowed, therefore it should not be done.
- The users must configure the correct IGP instance loopback IP addresses within the Rosen MVPN context and must configure the correct BGP policies (before Release 12.0.R4) for the feature to operate as expected. There is no verification that the address entered for the MVPN provider tunnel source address is such an address or is not a system IP address.

3.4.16 Inter-AS MVPN

The Inter-AS MVPN feature allows set-up of Multicast Distribution Trees (MDTs) that span multiple Autonomous Systems (ASes). This section focuses on multicast aspects of the Inter-AS MVPN solution.

To support Inter-AS option for MVPNs, a mechanism is required that allows setup of Inter-AS multicast tree across multiple ASes. Because of limited routing information across AS domains, it is not possible to setup the tree directly to the source PE. Inter-AS VPN Option A does not require anything specific to inter-AS support as customer instances terminate on ASBR and each customer instance is handed over to the other AS domain via a unique instance. This approach allows the users to provide full isolation of ASes, but the solution is the least scalable case, as customer instances across the network have to exist on ASBR.

Inter-AS MVPN Option B allows the users to improve upon the Option A scalability while still maintaining AS isolation, while Inter-AS MVPN Option C further improves Inter-AS scale solution but requires exchange of Inter-AS routing information and therefore is typically deployed when a common management exists across all ASes involved in the Inter-AS MVPN. The following sub-sections provide further details on Inter-AS Option B and Option C functionality.

3.4.16.1 BGP connector attribute

BGP connector attribute is a transitive attribute (unchanged by intermediate BGP speaker node) that is carried with VPNv4 advertisements. It specifies the address of source PE node that originated the VPNv4 advertisement.

With Inter-AS MVPN Option B, BGP next-hop is modified by local and remote ASBR during re-advertisement of VPNv4 routes. On BGP next-hop change, information about the originator of prefix is lost as the advertisement reaches the receiver PE node.

BGP connector attribute allows source PE address information to be available to receiver PE, so that a receiver PE is able to associate VPNv4 advertisement to the corresponding source PE.

3.4.16.2 PIM RPF vector

In case of Inter-AS MVPN Option B, routing information toward the source PE is not available in a remote AS domain, because IGP routes are not exchanged between ASes. Routers in an AS other than that of a source PE, have no routes available to reach the source PE and therefore PIM JOINS would never be sent upstream. To enable setup of MDT toward a source PE, BGP next-hop (ASBR) information from that PE's MDT-SAFI advertisement is used to fake a route to the PE. If the BGP next-hop is a PIM neighbor, the PIM JOINS would be sent upstream. Otherwise, the PIM JOINS would be sent to the immediate IGP next-hop (P) to reach the BGP next-hop. Because the IGP next-hop does not have a route to source PE, the PIM JOIN would not be propagated forward unless it carried extra information contained in RPF Vector.

In case of Inter-AS MVPN Option C, unicast routing information toward the source PE is available in a remote AS PEs/ASBRs as BGP 8277 tunnels, but unavailable at remote P routers. If the tunneled next-hop (ASBR) is a PIM neighbor, the PIM JOINS would be sent upstream. Otherwise, the PIM JOINS would be sent to the immediate IGP next-hop (P) to reach the tunneled next-hop. Because the IGP next-hop does not have a route to source PE, the PIM JOIN would not be propagated forward unless it carried extra information contained in RPF Vector.

To enable setup of MDT toward a source PE, PIM JOIN therefore carries BGP next hop information in addition to source PE IP address and RD for this MVPN. For option-B, both these pieces of information are derived from MDT-SAFI advertisement from the source PE. For option-C, both these pieces of information are obtained from the BGP tunneled route.

Use the following command to add the RPF vector to a PIM join at a PE router.

```
configure router pim rpfv
```

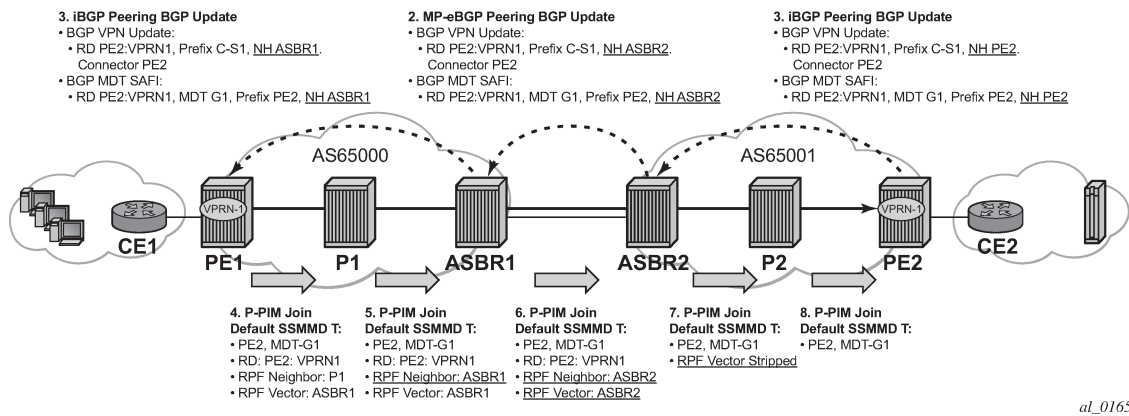
Also use the preceding command to configure P routers and ASBR routers to allow RPF Vector processing. If not configured, the RPF Vector is dropped and the PIM JOIN is processed as if the PIM Vector were not present.

For further details about RPF Vector processing please see [RFCs 5496, 5384 and 6513]

3.4.16.3 Inter-AS MVPN Option B

Inter-AS Option B is supported for Rosen MVPN PIM SSM using BGP MDT SAFI, PIM RPF Vector and BGP Connector attribute. [Figure 47: Inter-AS Option B default MDT setup](#) is an example of a default MDT setup.

Figure 47: Inter-AS Option B default MDT setup



SR OS Inter-AS Option B is designed to be standard compliant based on the following RFCs:

- | | |
|-----------------|--|
| RFC 5384 | The Protocol Independent Multicast (PIM) Join Attribute Format |
| RFC 5496 | The Reverse Path Forwarding (RPF) Vector TLV |
| RFC 6513 | Multicast in MPLS/BGP IP VPNs |

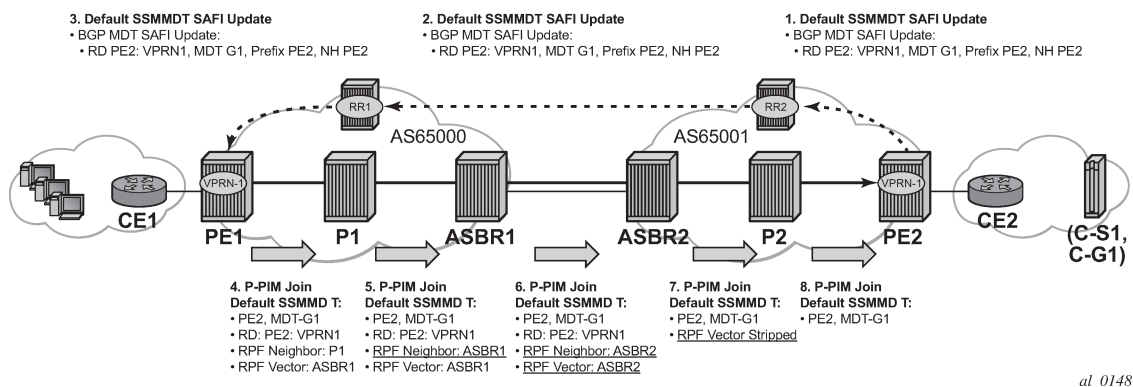
The SR OS implementation was designed also to interoperate with older routers Inter-AS implementations that do not comply with the RFC 5384 and RFC 5496.

3.4.16.4 Inter-AS MVPN Option C

Inter-AS Option C is supported for Rosen MVPN PIM SSM using BGP MDT SAFI and PIM RPF Vector.

Figure 48: Inter-AS Option C default MDT setup depicts a default MDT setup:

Figure 48: Inter-AS Option C default MDT setup



Additional restrictions for Inter-AS MVPN Option B and C support are the following:

- Inter-AS MVPN Option B is not supported with duplicate PE addresses.
- For Inter-AS Option C, BGP 8277 routes are installed into unicast rtm (rtable-u), unless routes are installed by some other means into multicast rtm (rtable-m), and Option C does not build core MDTs, therefore, rpf-table is configured to rtable-u or both.

Additional Cisco interoperability notes are the following:

RFC 5384	The Protocol Independent Multicast (PIM) Join Attribute Format
RFC 5496	The Reverse Path Forwarding (RPF) Vector TLV
RFC 6513	Multicast in MPLS/BGP IP VPNs

The SR OS implementation was designed to inter-operate with Cisco routers Inter-AS implementations that do not comply with the RFC5384 and RFC5496.

The following command configures RPF Vector processing for Inter-AS Rosen MVPN Option B and Option C.

```
configure router pim rpfv mvpn
```

For interoperability, when such a configuration is used, Cisco routers need to be configured to include RD in an RPF vector using the following command: **ip multicast vrf vrf-name rpf proxy rd vector**.

When Cisco routers are not configured to include RD in an RPF vector, use the following commands to configure the SR OS router (if supported):

- **MD-CLI**

```
configure router pim rpfv core
configure router pim rpfv mvpn
```

- **classic CLI**

```
configure router pim rpfv core mvpn
```

PIM joins received can be a mix of core and mvpn RPF vectors.

3.4.16.5 NG-MVPN non-segmented inter-AS solution

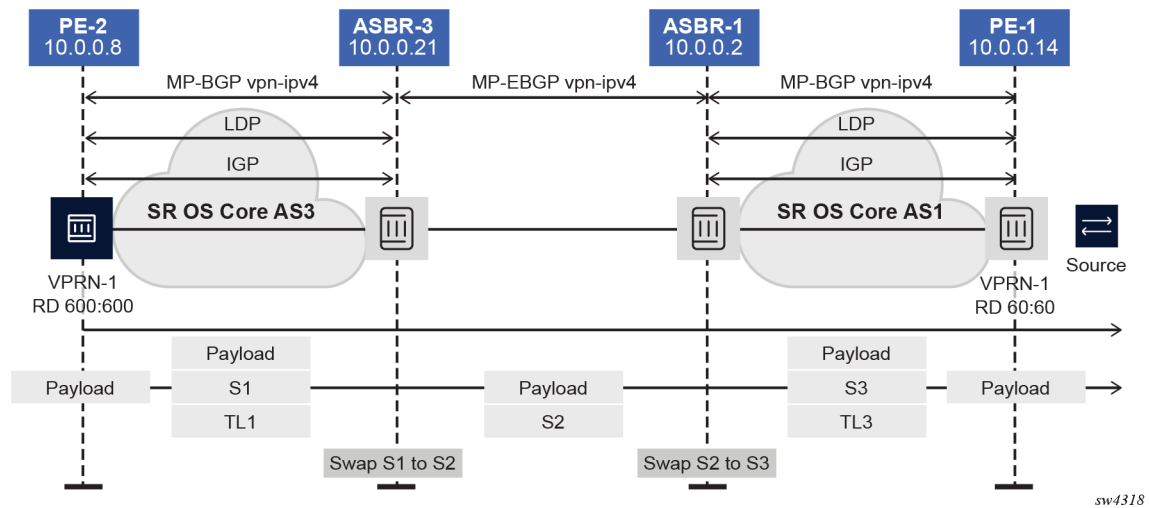
This feature allows multicast services to use segmented protocols and span them over multiple autonomous systems (ASs), as done in unicast services. As IP VPN or GRT services span multiple IGP areas or multiple ASs, either because of a network designed to deal with scale or as result of commercial acquisitions, the users may require Inter-AS VPN (unicast) connectivity. For example, an Inter-AS VPN can break the IGP, MPLS and BGP protocols into access segments and core segments, allowing higher scaling of protocols by segmenting them into their own islands. SR OS also allows for similar provision of multicast services and for spanning these services over multiple IGP areas or multiple ASs.

For multicast VPN (MVPN), SR OS previously supported Inter-AS Model A/B/C for Rosen MVPN; however, when MPLS was used, only Model A was supported for Next Generation Multicast VPN (NG-MVPN) and d-mLDP signaling.

For unicast VPRNs, the Inter-AS or Intra-AS Option B and C breaks the IGP, BGP and MPLS protocols at ABR routers (in case of multiple IGP areas) and ASBR routers (in case of multiple ASs). At ABR and ASBR routers, a stitching mechanism of MPLS transport is required to allow transition from one segment to next, as shown in [Figure 49: Unicast VPN Option B with segmented MPLS](#) and [Figure 50: Unicast VPN Option C with segmented MPLS](#).

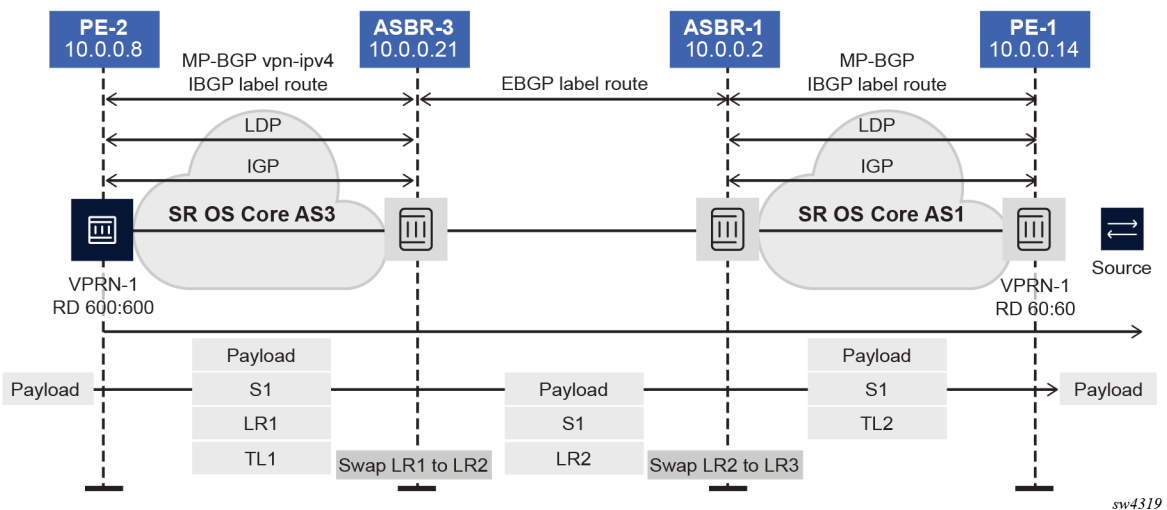
In [Figure 49: Unicast VPN Option B with segmented MPLS](#), the Service Label (S) is stitched at the ASBR routers.

Figure 49: Unicast VPN Option B with segmented MPLS



In Figure 50: Unicast VPN Option C with segmented MPLS, the 8277 BGP Label Route (LR) is stitched at ASBR1 and ASBR3. At ASBR1, the LR1 is stitched with LR2, and at ASBR3, the LR2 is stitched with TL2.

Figure 50: Unicast VPN Option C with segmented MPLS



Previously, in case of NG-MVPN, segmenting an LDP MPLS tunnel at ASBRs or ABRs was not possible. As such, RFC 6512 and 6513 used a non-segmented mechanism to transport the multicast data over P-tunnels end-to-end through ABR and ASBR routers. The signaling of LDP needed to be present and possible between two ABR routers or two ASBR routers in different ASs.



Note: For unicast VPNs, it was usually preferred to only have EBGP between ASBR routers. The non-segmented behavior of d-mLDP would have broken this by requiring LDP signaling between ASBR routers.

SR OS now has d-mLDP non-segmented intra-AS and inter-AS signaling for NG-MVPN and GRT multicast. The non-segmented solution for d-mLDP is possible for inter-ASs as Option B and C.

3.4.16.5.1 Non-segmented d-mLDP and inter-AS VPN

There are three types of VPN Inter-AS solutions:

- [Inter-AS Option A](#)
- [Inter-AS Option B](#)
- [Inter-AS option-C](#)

Options B and C use recursive opaque types 8 and 7 respectively, from [Table 8: Recursive opaque types](#).

Table 8: Recursive opaque types

Opaque type	Opaque name	RFC	SR OS use
1	Basic Type	RFC 6388	VPRN Local AS
3	Transit IPv4	RFC 6826	IPv4 multicast over mLDP in GRT
4	Transit IPv6	RFC 6826	IPv6 multicast over mLDP in GRT
7	Recursive Opaque (Basic Type)	RFC 6512	Inter-AS Option C MVPN over mLDP
8	Recursive Opaque (VPN Type)	RFC 6512	Inter-AS Option B MVPN over mLDP

3.4.16.5.1.1 Inter-AS Option A

In Inter-AS Option A, ASBRs communicate using VPN access interfaces, which need to be configured under PIM for the two ASBRs to exchange multicast information.

3.4.16.5.1.2 Inter-AS Option B

The recursive opaque type used for Inter-AS Option B is the Recursive Opaque (VPN Type), shown as opaque type 8 in [Table 8: Recursive opaque types](#).

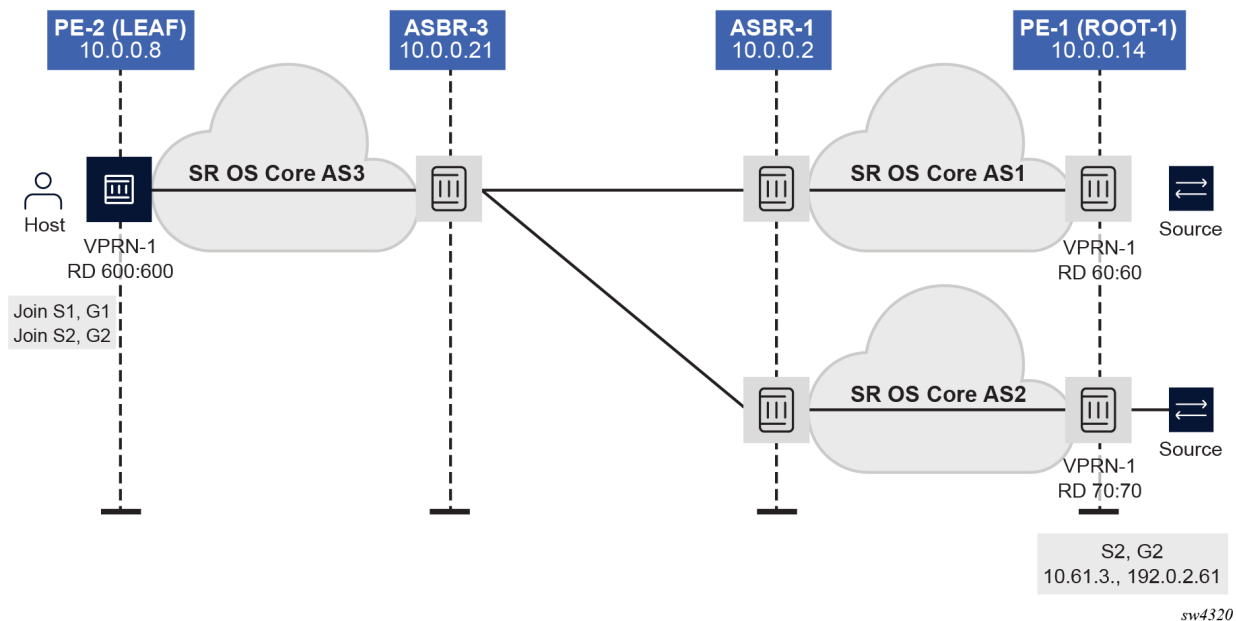
Inter-AS option B support for NG-MVPN

Inter-AS Option B requires additional processing on ASBR routers and recursive FEC encoding than that of Inter-AS Option A. Because BGP adjacency is not e2e, ASBRs must cache and use a PMSI route to build the tree. For that, mLDP recursive FEC must carry RD information—therefore, VPN recursive FEC is required (opaque type 8).

In Inter-AS Option B, the PEs in two different ASs do not have their system IP address in the RTM. As such, for NG-MVPN, a recursive opaque value in mLDP FEC is required to signal the LSP to the first ASBR in the local AS path.

Because the system IPs of the peer PEs (Root-1 and Root-2) are not installed on the local PE (leaf), it is possible to have two PEs in different ASs with same system IP address, as shown in [Figure 51: Identical system IP on multiple PEs \(Option B\)](#). However, SR OS does not support this topology. The system IP address of all nodes (root or leaf) in different ASs must be unique.

Figure 51: Identical system IP on multiple PEs (Option B)



For Inter-AS Option B and NG-MVPN, SR OS as a leaf does not support multiple roots in multiple ASs with the same system IP and different RDs; however, the first root that is advertised to an SR OS leaf is used by PIM to generate an mLDP tunnel to this actual root. Any dynamic behavior after this point, such as removal of the root and its replacement by a second root in a different AS, is not supported and the SR OS behavior is nondeterministic.

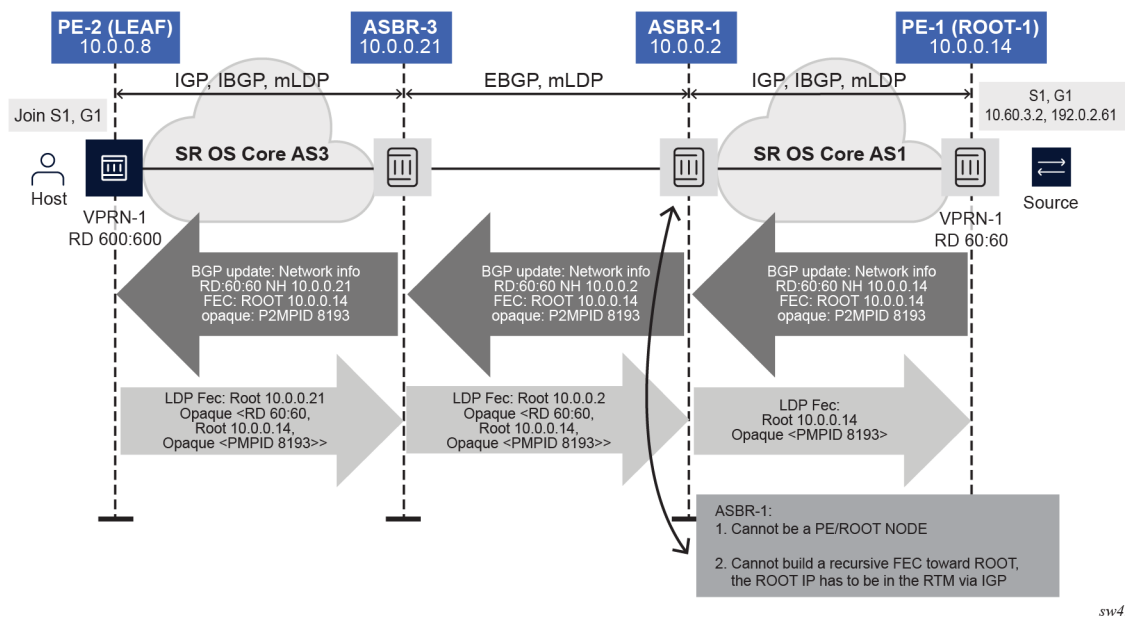
I-PMSI and S-PMSI establishment

I-PMSI and S-PMSI functionality follows RFC 6513 section 8.1.1 and RFC 6512 sections 3.1 and 3.2.1. For routing, the same rules as for GRT d-mLDP use case apply, but the VRR Route Import External community now encodes the VRF instance in the local administrator field.

Option B uses an outer opaque of type 8 and inter opaque of type 1 (see [Table 8: Recursive opaque types](#)).

[Figure 52: Non-segmented mLDP PMSI establishment \(Option B\)](#) depicts the processing required for I-PMSI and S-PMSI Inter-AS establishment.

Figure 52: Non-segmented mLDP PMSI establishment (Option B)



For non-segmented mLDP trees, A-D procedures follow those of the Intra-AS model, with the exception that NO EXPORT community must be excluded; LSP FEC includes mLDP VPN-recursive FEC.

For I-PMSI on Inter-AS Option B:

- A-D routes must be installed by ASBRs and next-hop information is changed as the routes are propagated, as shown in [Figure 52: Non-segmented mLDP PMSI establishment \(Option B\)](#).
- PMSI A-D routes are used to provide inter-domain connectivity on remote ASBRs.

On a receipt of an Intra-AS PMSI A-D route, PE2 resolves PE1's address (next-hop in PMSI route) to a labeled BGP route with a next-hop of ASBR3, because PE1 (Root-1) is not known via IGP. Because ASBR3 is not the originator of the PMSI route, PE2 sources an mLDP VPN recursive FEC with a root node of ASBR3, and an opaque value containing the information advertised by Root-1 (PE-1) in the PMSI A-D route, shown below, and forwards the FEC to ASBR 3 using IGP.

PE-2 LEAF FEC: (Root ASBR3, Opaque value {Root: ROOT-1, RD 60:60, Opaque Value: P2MPLSP-ID xx})

When the mLDP VPN-recursive FEC arrives at ASBR3, it notes that it is the identified root node, and that the opaque value is a VPN-recursive opaque value. Because Root-1 PE1 is not known via IGP, ASBR3 resolves the root node of the VPN-Recursive FEC using PMSI A-D (I or S) matching the information in the VPN-recursive FEC (the originator being PE1 (Root-1), RD being 60:60, and P2MP LSP ID xx). This yields ASBR1 as next hop. ASBR3 creates a new mLDP FEC element with a root node of ASBR1, and an opaque value being the received recursive opaque value, as shown below. ASBR then forwards the FEC using IGP.

ASBR-3 FEC: {Root ASBR 1, Opaque Value {Root: ROOT-1, RD 60:60, Opaque Value: P2MPLSP-ID xx}}

When the mLDP FEC arrives at ASBR1, it notes that it is the root node and that the opaque value is a VPN-recursive opaque value. As PE1's ROOT-1 address is known to ASBR1 through the IGP, no further recursion is required. Regular processing begins, using received Opaque mLDP FEC information.

ASBR-1 FEC: {Root: ROOT-1, Opaque Value: P2MP LSP-ID xx}



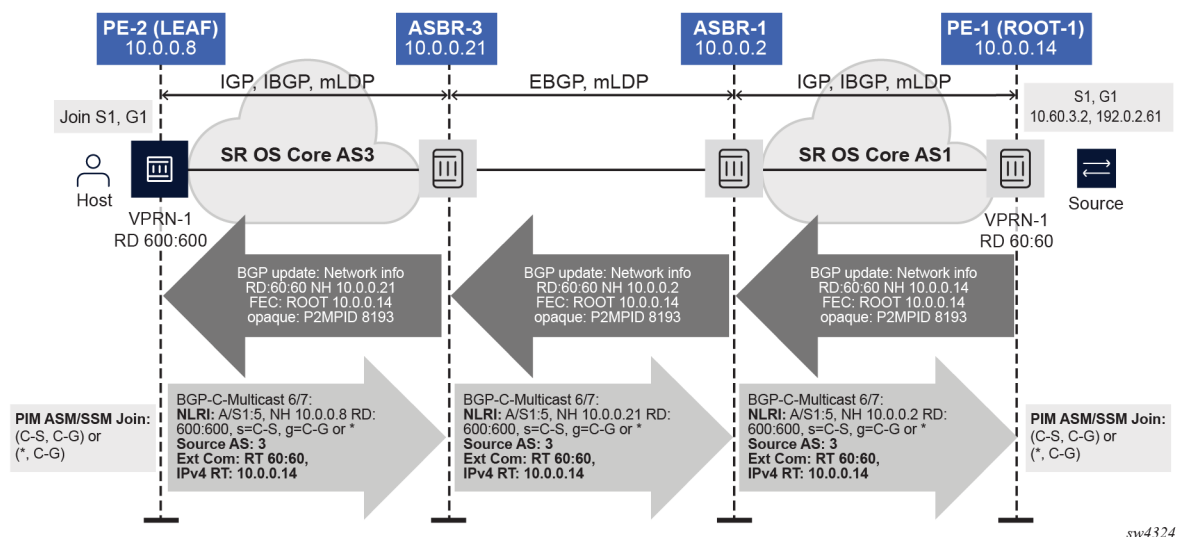
Note: VPN-Recursive FEC carries P2MPLSP ID. The P2MPLSP ID is used in addition to PE RD and Root to select a route to the mLDP root using the correct I-PMSI or S-PMSI route.

The functionality as described above for I-PMSI applies also to S-PMSI and (C-*, C-*) S-PMSI.

C-multicast Route Processing

C-multicast route processing functionality follows RFC 6513 section 8.1.2 (BGP used for route exchange). The processing is analogous to BGP Unicast VPN route exchange described in [Figure 49: Unicast VPN Option B with segmented MPLS](#) and [Figure 50: Unicast VPN Option C with segmented MPLS](#). [Figure 53: Non-segmented mLDP C-multicast exchange \(Option B\)](#) shows C-multicast route processing with non-segmented mLDP PMSI details.

Figure 53: Non-segmented mLDP C-multicast exchange (Option B)

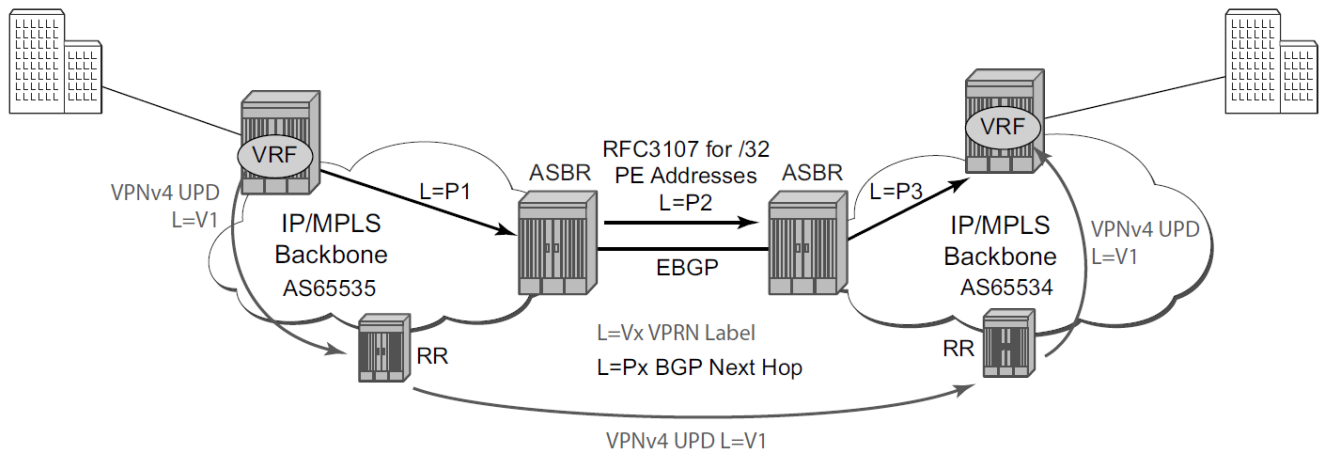


sw4324

3.4.16.5.1.3 Inter-AS option-C

Inter-AS option-C allows a higher scale of VPRNs across AS boundaries and expands the trust model between the ASNs. As a result, this model is typically used within a single organization that may have multiple ASNs. The following figure shows inter-AS option-C.

Figure 54: Inter-AS option-C



OSSG257

The inter-AS option-C model differs from option-B, in that all direct knowledge of the remote AS is contained in and limited to the ASBR in an option-B network. The ASBR performs all necessary mapping functions, and the PE routers do not need to perform additional functions other than those performed by a non-inter-AS VPRN.

In option-C, however, knowledge from the remote AS is distributed throughout the local AS. This distribution allows for higher scalability, but also requires that all PEs and ASBRs involved in the inter-AS VPRNs participate in the exchange of inter-AS routing information. The ASBRs distribute reachability information for remote PE system IP addresses only by exchanging MP-eBGP labeled routes, as defined in RFC 8277, *Using BGP to Bind MPLS Labels to Address Prefixes*.

In option-C, the ASBRs distribute reachability information for remote PE's system IP addresses only. This is done between the ASBRs by exchanging MP-EBGP labeled routes, using RFC 8277, *Using BGP to Bind MPLS Labels to Address Prefixes*. Either RSVP-TE or LDP LSP can be selected to resolve next-hop for multihop EBGP peering by using the following command.

```
configure router bgp next-hop-resolution labeled-routes transport-tunnel
```

Distribution of VPRN routing information is handled by either direct MP-BGP peering between PEs in the different ASNs, or by one or more route reflectors in the ASN.

3.4.16.5.1.3.1 Inter-AS Option C support for NG-MVPN

For Inter-AS Option C, on a leaf PE, a route exists to reach root PE's system IP and, as ASBRs can use BGP unicast routes, recursive FEC processing using BGP unicast routes, and not VPN recursive FEC processing using PMSI routes, is required.

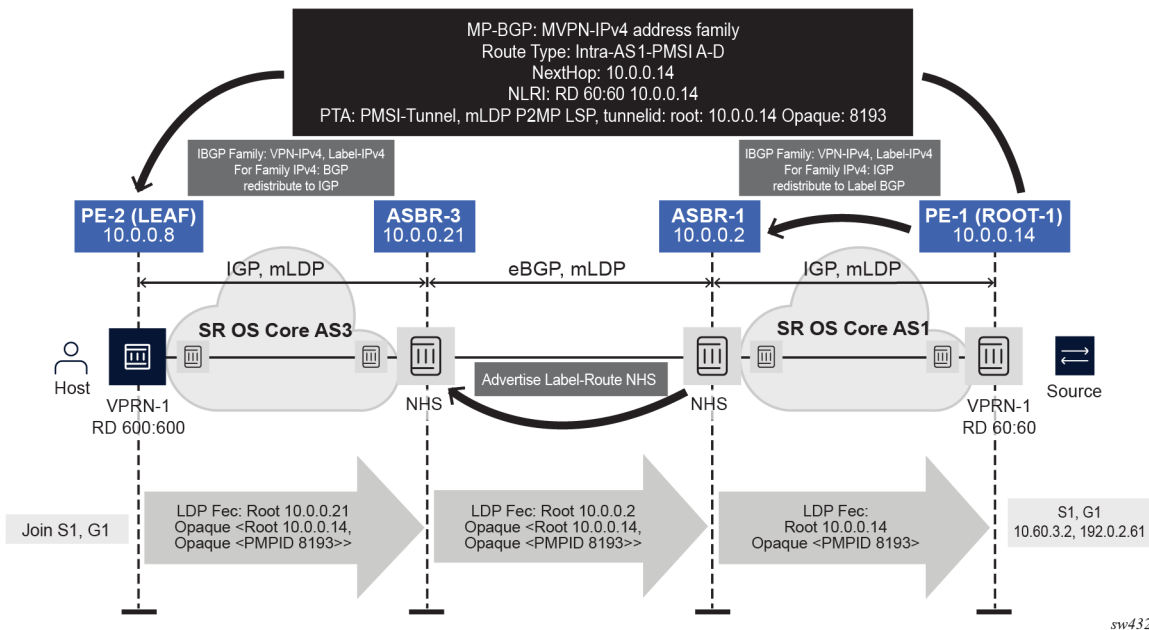
I-PMSI and S-PMSI establishment

I-PMSI and S-PMSI functionality follows RFC 6513 section 8.1.1 and RFC 6512 Section 2. The same rules as per the GRT d-mLDP use case apply, but the VRR Route Import External community now encodes the VRF instance in the local administrator field.

Option C uses an outer opaque of type 7 and inter opaque of type 1.

Figure 55: Non-segmented mLDP PMSI establishment (Option C) shows the processing required for I-PMSI and S-PMSI Inter-AS establishment.

Figure 55: Non-segmented mLDP PMSI establishment (Option C)



For non-segmented mLDP trees, A-D procedures follow those of the Intra-AS model, with the exception that NO EXPORT Community must be excluded; LSP FEC includes mLDP recursive FEC (and not VPN recursive FEC).

For I-PMSI on Inter-AS Option C:

- A-D routes are not installed by ASBRs and next-hop information is not changed in MVPN A-D routes.
- BGP-labeled routes are used to provide inter-domain connectivity on remote ASBRs.

On a receipt of an Intra-AS I-PMSI A-D route, PE2 resolves PE1's address (N-H in PMSI route) to a labeled BGP route with a next-hop of ASBR3, because PE1 is not known via IGP. PE2 sources an mLDP FEC with a root node of ASBR3, and an opaque value, shown below, containing the information advertised by PE1 in the I-PMSI A-D route.

PE-2 LEAF FEC: {root = ASBR3, opaque value: {Root: ROOT-1, opaque value: P2MP-ID xx}}

When the mLDP FEC arrives at ASBR3, it notes that it is the identified root node, and that the opaque value is a recursive opaque value. ASBR3 resolves the root node of the Recursive FEC (ROOT-1) to a labeled BGP route with the next-hop of ASBR1, because PE-1 is not known via IGP. ASBR3 creates a new mLDP FEC element with a root node of ASBR1, and an opaque value being the received recursive opaque value.

ASBR3 FEC: {root: ASBR1, opaque value: {root: ROOT-1, opaque value: P2MP-ID xx}}

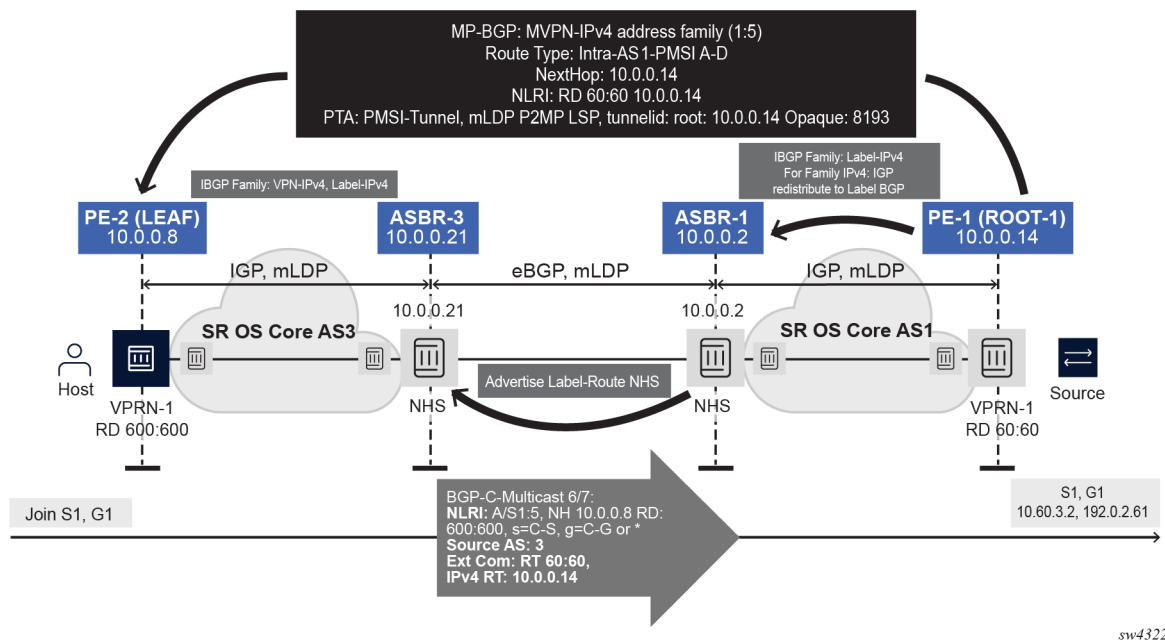
When the mLDP FEC arrives at ASBR1, it notes that it is the root node and that the opaque value is a recursive opaque value. As PE-1's address is known to ASBR1 through the IGP, no further recursion is required. Regular processing begins, using the received Opaque mLDP FEC information.

The functionality as described above for I-PMSI applies to S-PMSI and (C-*, C-*) S-PMSI.

C-multicast route processing

C-multicast route processing functionality follows RFC 6513 section 8.1.2 (BGP used for route exchange). The processing is analogous to BGP Unicast VPN route exchange. [Figure 56: Non-segmented mLDP C-multicast exchange \(Option C\)](#) shows C-multicast route processing with non-segmented mLDP PMSI details.

Figure 56: Non-segmented mLDP C-multicast exchange (Option C)



LEAF node cavities



Caution: The SR OS ASBR does not currently support receiving a non-recursive opaque FEC (opaque type 1).

The LEAF (PE-2) has to have the ROOT-1 system IP installed in RTM via BGP. If the ROOT-1 is installed in RTM via IGP, the LEAF does not generate the recursive opaque FEC. As such, the ASBR 3 does not process the LDP FEC correctly.

3.4.16.5.2 Configuration example

No configuration is required for Option B or Option C on ASBRs. For Option B, use the following command to configure inter-as-non-segmented MLDP through the ASBR router:

- **MD-CLI**

```
configure router bgp inter-as-vpn
```

- **classic CLI**

```
configure router bgp enable-inter-as-vpn
```

A policy is required for a root or leaf PE for removing the NO_EXPORT community from MVPN routes, which can be configured using an export policy on the PE.

The following example displays a policy configured on PEs to remove the **no-export** community.

Example: MD-CLI

```
[ex:/configure router "Base" bgp]
A:admin@node-2# info
  community "no-export" {
    member "no-export" { }
  }
  policy-statement "remNoExport" {
    default-action {
      action-type accept
      community {
        remove ["no-export"]
      }
    }
  }
}
```

Example: classic CLI

```
A:node-2>config>router>policy-options# info
-----
  community "no-export"
    members "no-export"
  exit
  policy-statement "remNoExport"
    default-action accept
      community remove "no-export"
    exit
  exit
```

The following is an example for configuring in BGP the policy in a global, group, or peer context.

Example: MD-CLI

```
[ex:/configure router "Base" bgp]
A:admin@node-2# info
  vpn-apply-export true
  export {
    policy ["remNoExport"]
  }
```

Example: classic CLI

```
A:node-2>config>router>bgp# info
-----
  vpn-apply-export
  export "remNoExport"
```

3.4.16.5.3 Inter-AS non-segmented MLDP

See the "Inter-AS Non-segmented MLDP" section of the *7705 SAR Gen 2 MPLS Guide* for more information.

3.4.16.5.4 ECMP

See the "ECMP" section of the *7705 SAR Gen 2 MPLS Guide* for more information about ECMP.

3.4.17 mLDP non-Segmented intra-AS (inter-area) MVPN solution

SR OS supports intra-AS (Inter-Area) option B and C. The following interaction between inter and intra is as follows:

- Intra-AS option B with inter-AS option B
- Intra-AS option C with inter-AS option C

3.4.17.1 Intra-AS and Inter-AS Option B

For Intra/Inter-as option B, the root is not visible on the leaf. LDP is responsible for building the recursive FEC and signaling the FEC to ABR/ASBR on the leaf. ABR/ASBR must have the PMSI AD router to rebuild the FEC (recursive or basic) depending on whether they are connected to another ABR/ASBR or root node. As such, LDP must import the MVPN PMSI AD routes. To save resources, importing MVPN PMSI AD routes are performed manually by the user using configuration commands.

Use the following command to configure LDP to request BGP to provide the LDP task with all the MVPN PMSI AD routes and LDP internally caches these routes.

```
configure router ldp import-pmsi-routes mvpn-no-export-community
```

When **mvpn-no-export-community** is disabled, MVPN discards the catch routes to save resources.

In a scenario where a node running an older image that does not support the **mvpn-no-export-community** command and that node upgrades to a new image that does support the **mvpn-no-export-community** command, and if the older image had an inter-as MVPN configuration, after the upgrade to the newer image, the **mvpn-no-export-community** command is enabled by default, to ensure a smooth upgrade. This is to verify that all the routes are imported to mLDP, so the inter-as functionality works after the upgrade.

SR OS supports two major upgrades to enable the MVPN command if an upgrade to a load supporting this knob.

3.4.17.2 MVPN next hop self on ABRs

For option B, the ABR routers must change the next-hop of MVPN AD routes to be the ABR systemIP or the loopback IP for core diversity. Currently, the **next-hop-self** BGP command does not change the next hop of the MVPN AD routes. This functionality will be available in a future release.

In the meantime, a BGP policy can be used to change the MVPN AD routes next hop at the ABR.

3.4.17.2.1 MVPN next-hop-self policy example

MVPN Type 1 route (intra-AS IPMSI AD route) and MVPN Type 3 (S-PMSI AD route) must have a policy to set their next hop to be the ABR systemIP. In the following example, the ABR systemIP is 10.20.1.4 with the same token as the unicast vpn-ipv4 family and can be configured within the policy to have the next hop changed to the ABR systemIP.

Configure three policies on all ABRs:

- a policy to change mvpn-ipv4 IntraAD Route Type 1 next hop to next-hop-self
- a policy to change vpn-ipv4 next hop to next-hop-self
- a policy to change mvpn-ipv4 IntraAD Route Type 3 to next-hop-self

Example: MD-CLI

```
[ex:/configure policy-options]
A:admin@node-2# info

policy-statement "mod_nh_10.20.1.4" {
  entry 1 {
    from {
      mvpn-type intra-as-ipmsi-auto-discovery
    }
    to {
      neighbor {
        ip-address 10.20.1.4
      }
    }
    action {
      action-type accept
    }
  }
  default-action {
    action-type next-policy
  }
}
policy-statement "mod_nh_spmsi_10.20.1.4" {
  entry 1 {
    from {
      mvpn-type s-pmsi-auto-discovery
    }
    to {
      neighbor {
        ip-address 10.20.1.4
      }
    }
    action {
      action-type accept
    }
  }
  default-action {
    action-type next-policy
  }
}
policy-statement "mod_nh_vpn_10.20.1.4" {
  entry 1 {
    from {
      family [vpn-ipv4]
    }
    to {
      neighbor {
```

```

        ip-address 10.20.1.4
    }
    action {
        action-type accept
    }
}
default-action {
    action-type next-policy
}
}

```

Example: classic CLI

```

A:node-2>config>router>policy-options# info
-----
    policy-statement "mod_nh_10.20.1.4"
        entry 1
            from
                mvpn-type 1
            exit
            action accept
                next-hop 10.20.1.4
            exit
        exit
        default-action next-policy
        exit
    exit
    policy-statement "mod_nh_vpn_10.20.1.4"
        entry 1
            from
                family vpn-ipv4
            exit
            action accept
                next-hop 10.20.1.4
            exit
        exit
        default-action next-policy
        exit
    exit
    policy-statement "mod_nh_spmsi_10.20.1.4"
        entry 1
            from
                mvpn-type 3
            exit
            action accept
                next-hop 10.20.1.4
            exit
        exit
        default-action next-policy
        exit
    exit
-----

```

3.4.17.2.2 LDP configuration example

Use the following commands to import all inter-AS and intra-AS (inter-area) routes on all ABR and non-ABR routers.

```
configure router ldp import-pmsi-routes mvpn
configure router ldp import-pmsi-routes mvpn-no-export-community
```

The following example shows the configuration.

Example: MD-CLI

```
[ex:/configure router "Base" ldp]
A:admin@node-2# info
    mp-mbb-time 10
    generate-basic-fec-only true
    fast-reroute {
    }
    import-pmsi-routes {
        mvpn true
        mvpn-no-export-community true
    }
```

Example: classic CLI

```
A:node-2>config>router>ldp# info
-----
    fast-reroute
    mp-mbb-time 10
    generate-basic-fec-only
    import-pmsi-routes
        mvpn
        mvpn-no-export-community
    exit
```

3.4.17.2.3 BGP configuration example

Use the following commands to import on ABR routers the MVPN AD **next-hop-self** policies and the **vpn-ipv4** family.

```
configure router bgp group next-hop-self
configure router bgp group family vpn-ipv4
```

In addition, for unicast **vpn-ipv4** connectivity, use the following command to configure inter-as-non-segmented MLDP through the ASBR router:

- **MD-CLI**

```
configure router bgp inter-as-vpn
```

- **classic CLI**

```
configure router bgp enable-inter-as-vpn
```

The following example shows the BGP configuration.

Example: MD-CLI

```
[ex:/configure router "Base" bgp]
A:admin@node-2# info
  connect-retry 10
  keepalive 10
  vpn-apply-export true
  vpn-apply-import true
  inter-as-vpn true
  hold-time {
    seconds 30
  }
  family {
    vpn-ipv4 true
    ipv6 true
    vpn-ipv6 true
    mvpn-ipv4 true
    mcast-vpn-ipv4 true
    mvpn-ipv6 true
    mcast-vpn-ipv6 true
  }
  ...
  rapid-update {
    vpn-ipv4 true
    vpn-ipv6 true
    mvpn-ipv4 true
    mcast-vpn-ipv4 true
    mvpn-ipv6 true
    mcast-vpn-ipv6 true
  }
  group "ibgp_A" {
    next-hop-self true
    cluster {
      cluster-id 10.20.1.2
    }
  }
  group "ibgp_D" {
    next-hop-self true
    cluster {
      cluster-id 10.180.4.2
    }
  }
}
```

Example: classic CLI

```
A:node-2>config>router>bgp# info
-----
family ipv4 ipv6 vpn-ipv4 vpn-ipv6 mvpn-ipv4 mcast-vpn-ipv4 mvpn-ipv6 mcast-
vpn-ipv6
  vpn-apply-import
  vpn-apply-export
  connect-retry 10
  keepalive 10
  hold-time 30
  enable-inter-as-vpn
  rapid-update vpn-ipv4 vpn-ipv6 mvpn-ipv4 mcast-vpn-ipv4 mvpn-ipv6 mcast-vpn-
ipv6
  group "ibgp_A"
    next-hop-self
    cluster 10.20.1.2
    export "mod_nh_10.20.1.2" "mod_nh_spmsi_10.20.1.2" "mod_nh_vpn_10.20.1.2"
    neighbor 10.20.1.1
    local-address 10.20.1.2
```

```

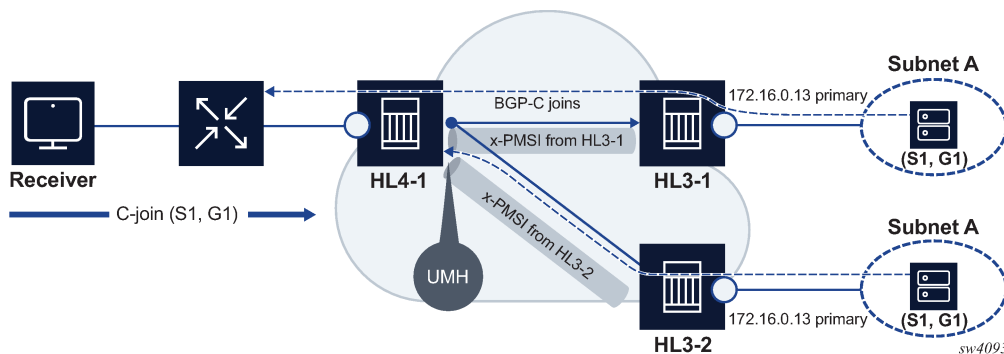
        med-out 100
        peer-as 100
    exit
exit
group "ibgp_D"
    next-hop-self
    cluster 10.180.4.2
    export "mod_nh_10.20.1.2" "mod_nh_spmsi_10.20.1.2" "mod_nh_vpn_10.20.1.2"
exit
no shutdown
-----

```

3.4.18 UMH redundancy for bandwidth monitoring using a single IOM

Bandwidth monitoring is used in MVPN for NG-MVPN and mLDP transport. It is used for multicast source redundancy, where both sources have the same IP address but are connected to two different root nodes. Bandwidth monitoring can be used with basic or recursive mLDP FEC. Upstream Multicast Hop (UMH) redundancy for bandwidth monitoring is supported for mLDP basic FEC and recursive FEC type 7 and 8 only.

Figure 57: Bandwidth monitoring



With bandwidth monitoring, the leaf node sends a single (S1,G1) join to both root nodes. PIM SSM and ASM can be used between the receiver and the leaf, or between the UMH and the source. For ASM, bandwidth monitoring works only when traffic is switched from $\langle *, G \rangle$ to $\langle S, G \rangle$.

After the source starts the multicast flow toward the root PEs, both root nodes transport the traffic to the leaf node on the PMSI (I-PMSI or S-PMSI).

The leaf listens to the active PMSI, blocks the other PMSI, and monitors the traffic rate on both the active and inactive PMSI. For faster than 50 ms switchover, both the active and the inactive PMSIs must arrive on the same IOM, because a single IOM must make the decision about which PMSI the leaf listens to and which PMSI to block.

The threshold for the rate of traffic lost between the active PMSI and the inactive PMSI is configured on the leaf PE. If the rate exceeds the configured value, the traffic switches from the active PMSI to the inactive PMSI. Rate monitoring is per PMSI, and not per (C-S,C-G).

After the active PMSI traffic rate is stored, there is a revertive behavior, which has a configurable timer. The revertive timer starts after the active PMSI traffic is recovered. When the timer expires and the primary PMSI traffic is stable, the traffic is switched back to the primary path. If the traffic goes below the threshold while the timer is decrementing, the timer is reset. This feature supports 1K of PMSI switchovers within 50 ms.

3.4.18.1 Fault recovery mitigation at PMSI switchover time

<S,G> switching between I-PMSI and S-PMSI is not symmetrical (synchronized in time) on the active and the inactive UMH. While the active UMH attempts to switch an <S,G> between I-PMSI and S-PMSI, the active PMSI traffic rate arriving from the active UMH may be different from that arriving from the inactive UMH. This asymmetrical behavior can generate a premature switch from the active PMSI to the inactive PMSI.

The traffic rate delta can be set to account for this behavior. For example, if a 1080P channel uses 5 Mb/s, the traffic rate delta can be set to 15 Mb/s to avoid the switchover from the primary to the secondary PMSI if one or two 1080 <S,G>s are switched between I-PMSI and S-PMSI. This provides a 10 Mb/s tolerance of asymmetric traffic.

3.4.18.2 S-PMSI behavior

If the network FDV is large or the sources are not synchronized, switching from I-PMSI to S-PMSI can happen at a different time on the primary and backup UMHs. This can cause asymmetric traffic on the I-PMSI and S-PMSI, resulting in a switch from the active UMH. The <S,G> traffic can arrive for the I-PMSI from the backup UMH and for S-PMSI from the active UMH, which causes temporary duplicate traffic until both UMHs switch to S-PMSI.

Multistream S-PMSI provides a solution for this case by mapping an <S,G> to an S-PMSI. The <S,G> is locked to the multistream S-PMSI, which is always configured and never torn down, even if the traffic goes down to 0, so the multistream S-PMSI is not susceptible to S-PMSI traffic drops.

Use the following commands to configure the maximum number of S-PMSI for the MVPN selective provider tunnel.

- **MD-CLI**

```
configure service vprn mvpn provider-tunnel selective mldp maximum-p2mp-spmsi
configure service vprn mvpn provider-tunnel selective rsvp maximum-p2mp-spmsi
```

- **classic CLI**

```
configure service vprn mvpn provider-tunnel selective maximum-p2mp-spmsi
```

The number of <S,G>s must be less than, or equal to, this number. Otherwise, different <S,G>s can switch to the S-PMSI and when the S-PMSI limit is exhausted, the primary and backup UMHs become out of sync.

3.4.18.3 Bandwidth monitoring on single IOMs

Bandwidth monitoring is supported on single IOMs, that is, both the active and the backup PMSI terminate on the same IOM. The IOM monitors the statistics of both PMSIs and makes the switchover decision. The IOM does not include the CPM in any of the bandwidth monitoring decisions, which ensures fast detection times and switchover times under 50 ms.

All leaf and bud nodes must be configured with the same UMH PEs, I-PMSI and S-PMSI, and bandwidth threshold configuration to avoid traffic drops.

All LAG members must be in the same IOM that is performing the bandwidth monitoring function. The LAG interfaces spanning between multiple port members that belong to different IOMs have an unpredictable behavior, including traffic duplication.

3.4.18.4 ASM behavior

Bandwidth monitoring is supported with ASM only after the traffic is switched from <*,G> to <S,G>. Traffic arriving on separate IOMs from the active UMH and the inactive UMH results in traffic duplication because the pairing of the active and inactive UMH is per IOM, and the IOMs do not have a view of the pair. If the active traffic and the backup traffic arrive on different IOMs, each IOM treats the flow as the active flow and processes the traffic accordingly.

3.4.18.5 Low traffic rate

At low traffic rates, the packet on the active PMSI and the packet on the inactive PMSI can arrive at different times. If the packets arrive at the point that the statistics are read, there can be an inconsistency, resulting in a switchover. To avoid a switchover, UMH redundancy using bandwidth monitoring should be used only when the traffic rate is higher than 2 or 3 packets per second.

3.4.18.6 Revertive timer

Use the following commands to configure the revertive timer.

```
configure service vprn mvpn provider-tunnel inclusive umh-rate-monitoring revertive-timer
configure service vprn mvpn provider-tunnel selective umh-rate-monitoring group source
revertive-timer
```

If the active PMSI traffic goes below the threshold while the revertive timer is running, the timer is reset.

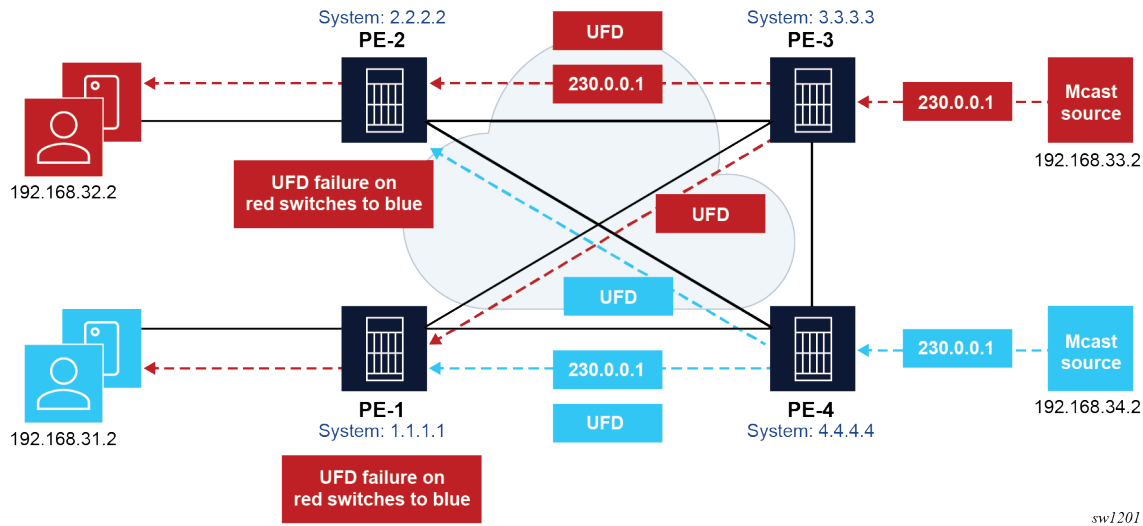
3.4.18.7 MVPN upstream PE fast failover

MVPN upstream UMH fast failover is supported for P2MP SR PMSI, as illustrated in [Figure 58: Example of MVPN upstream PE fast failover](#). The upstream PE failure detection uses the Unidirectional Forwarding Detection (UFD) method. A downstream PE (receiver PE) supports fast upstream failover using the capability:

- to select two UMH nodes
- to monitor the upstream PE health using UFD

The tunnel status is monitored using the UFD session status received over the P-tunnel; the C-flow source is declared to be active based on this monitoring information. If multiple nodes are sourcing C-flows, the receiver PE node can choose to receive traffic from a primary and a standby source, but forwards only the multicast stream received from the primary source. If the detected P-tunnel status is down, the multicast receiver PE forwards the traffic received over the standby P-tunnel.

Figure 58: Example of MVPN upstream PE fast failover



3.4.18.7.1 MVPN upstream PE fast failover for tree SID

MVPN upstream PE fast failover for tree SID is supported as follows:

- Only inclusive PMSIs are supported.
- SM and SSM modes are supported. For SM, only fast switchover is supported on SPT. Fast protection is not supported on a shared tree.
- UFD sessions with 10-millisecond interval are supported on the CPM.
- For the MVPN UMH redundancy feature with BFD fault detection, if using P2MP RSVP or P2MP SR-policy transport, the switchover time from the primary to the standby tunnels increases linearly as the number of tunnel pairs increases.
- Traffic duplication for restoration of the primary stream from the standby can occur for up to a second or more, depending on the CPU load during switchover.
- The P2MP policy tunnel ID is used as the UFD discriminator. Consequently, the solution is not interoperable with other vendors.

3.4.18.8 Multicast-only Fast Reroute

Multicast-only Fast Reroute (MoFRR) is not supported when UMH redundancy with bandwidth monitoring is enabled.

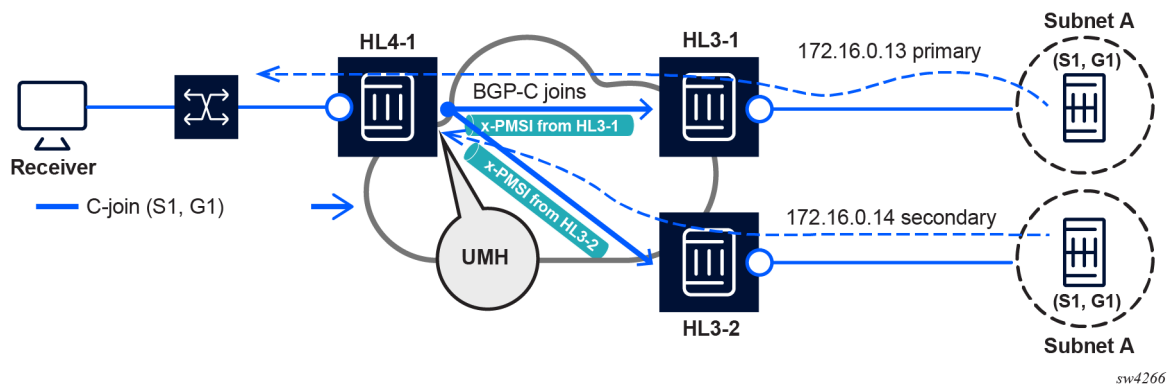
3.4.19 UMH redundancy for bandwidth monitoring using multiple IOMs

Upstream Multicast Hop (UMH) redundancy for bandwidth monitoring using a single IOM is effective if the active and backup P-Multicast Service Interface (PMSI) traffic arrives on a single IOM. However, if multicast streams from the active PMSI and backup PMSI arrive on different IOMs, it is not possible to collect and compare statistics between the two different IOMs.

An alternative solution is to use two IOMs for UMH redundancy for bandwidth monitoring. This approach uses PXC ports to forward the multicast stream from the active PMSI and backup PMSI to a single monitoring PMSI. The monitoring PMSI compares the multicast streams from a bandwidth perspective and decides which stream is forwarded to the host.

The following figure shows a UMH redundancy deployment for bandwidth monitoring using multiple IOMs.

Figure 59: UMH redundancy for bandwidth monitoring using multiple IOMs



Implementation considerations

Consider the following guidelines when implementing UMH redundancy with two IOMs using PXC ports:

- The monitoring IOM can be one of the IOMs on which the active or backup PMSI traffic arrives, or it can be a separate IOM.
- If you implement the UMH redundancy solution with two IOMs using PXC ports, you must also use PXC to implement the single IOM solution.



Note: The single IOM solution can monitor two PMSIs on the ingress IOM without using PXC. However, if you enable the solution with multiple IOMs using PXC, you must migrate the single IOM solution to the PXC solution. That is, the non-PXC solution with single-IOM monitoring and the PXC solution with two-IOM monitoring cannot coexist.

See [Bandwidth monitoring on single IOMs](#) for more information about the single IOM solution.

3.4.19.1 Implementation overview

Implementation and configuration guidelines

The following provides an overview of the Nokia implementation of UHM redundancy for bandwidth monitoring using two IOMs:

- When bandwidth monitoring using two IOMs is configured using PXC, the system cannot function in single-IOM (none PXC mode) bandwidth monitoring mode. You must implement the single-IOM bandwidth monitoring using PXC solution.
- An IOM is required to perform monitoring. The monitoring IOM (on which the PMSI traffic arrives) can be the active or standby ingress UMH IOM, or it can be a separate IOM. The PXC solution is required in either case, and the traffic to the PXC is forwarded over the fabric, consuming bandwidth.

- A PXC port must be created on the monitoring IOMs, either on the physical port or on the internal Ethernet port.
- The PMSI traffic from the active or standby UMH IOMs is forwarded to the PXC port, where the bandwidths of the multicast streams are monitored and the decision is made about which PMSI traffic to forward to the host.
 - The router uses the ILM lookups on the ingress UMH IOMs to forward the PMSI traffic to the monitoring IOM. These ILM lookups must be programmed on the system. For example, the system must configure an ILM entry for the incoming label (IL) of the PMSIs that are used for bandwidth monitoring. This entry is swapped with a Local System-assigned Label (LSL) with an NHLFE that points to the PXC interface on the monitoring IOM. The ingress UMH IOM uses this ILM entry to forward the PMSI streams to the monitoring IOM.
 - On the monitoring IOM, the system must create another set of ILM entries for the LSL to pop the label and perform forwarding to the egress ports upon which the hosts reside.
- To implement the comparison of the PMSI multicast traffic, the monitoring IOM is configured with a delta bandwidth in megabytes (MB). If the delta bandwidth between the PMSI traffic rate of the active UHM IOM is lower than the PMSI traffic rate of the inactive UHM IOM, as compared to the rate dictated by the configured delta bandwidth, the monitoring IOM blocks the active PMSI and sends the standby PMSI multicast stream to the receivers.

To support this functionality, the router uses a type of MVPN UMH rate monitoring Forwarding Path Extension (FPE). Use the following commands to specify an FPE for UMH rate monitoring.

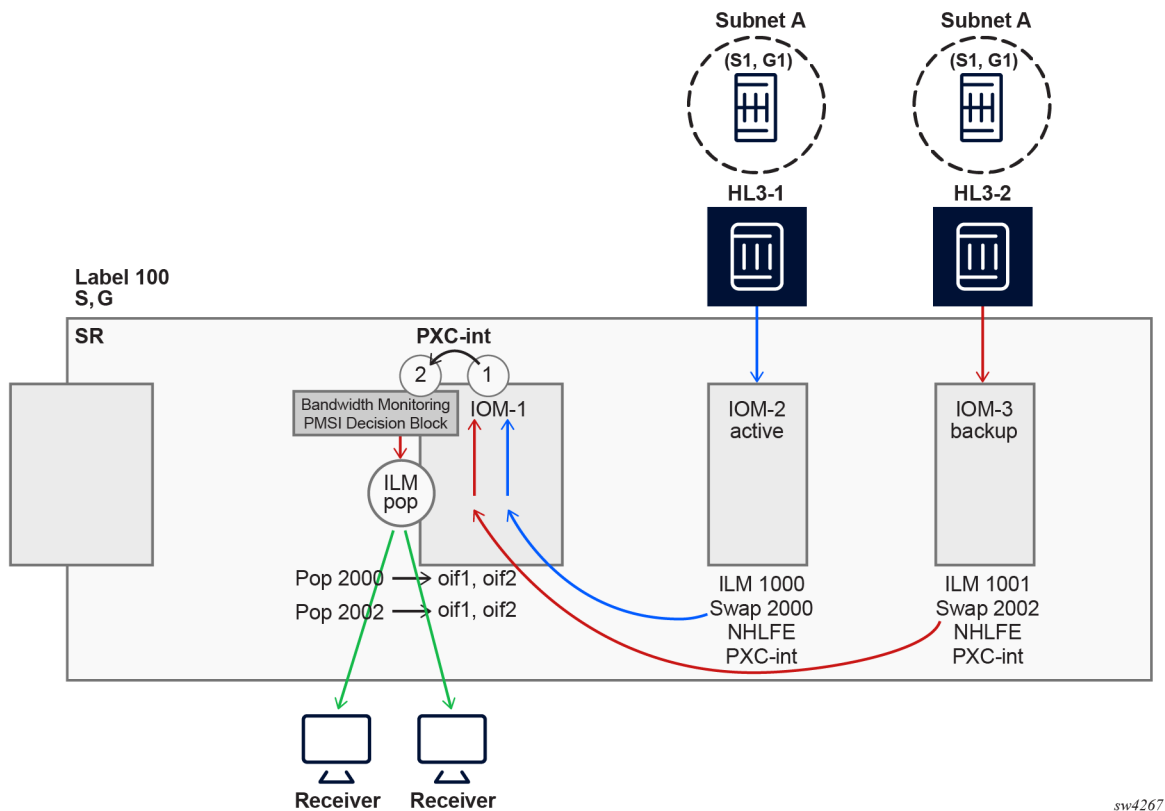
```
configure service vprn mvpn provider-tunnel inclusive umh-rate-monitoring fpe
configure service vprn mvpn provider-tunnel selective umh-rate-monitoring group source fpe
```



Note: If you use the UMH redundancy solution with two IOMs using PXC ports, you must also use PXC to implement the single IOM solution.

The following figure shows the configuration of bandwidth monitoring to a dedicated IOM.

Figure 60: Bandwidth monitoring to a dedicated IOM



Monitoring using redundant IOMs

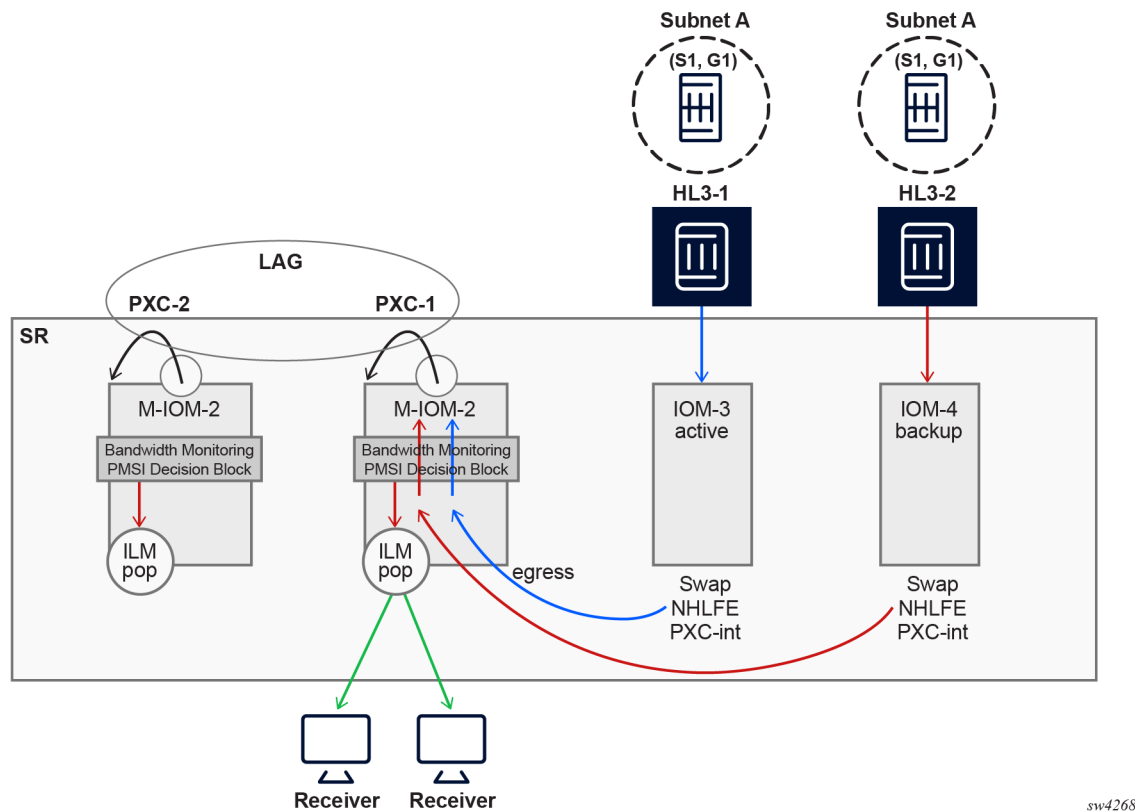
The two monitoring IOMs used for monitoring redundancy can be the same as the ingress active or backup IOMs, or a separate IOM. A PXC port can be created on each IOM and added as a LAG member to the LAG bundle using the **per-link-hash** command. An FPE can have many types. For UMH redundancy, **mvpn-umh-rate-monitoring** is configured as the FPE type, and is referenced using the **fwd-path-ext** command to create the interface on this LAG bundle with its own Martian IP address. When received at the ingress UMH IOMs, the router forwards the PMSI traffic to the active LAG IOM or, if there is a failure on the active LAG IOM, the router forwards the PMSI traffic to the standby IOM.



Note: Bandwidth monitoring uses the per-link hashing mode for the LAG configuration. In per-link hashing mode, the user cannot specify the active and standby member. It is therefore not possible to predict which PXC port is the active port and which is the standby. One of the IOMs in the redundant pair randomly becomes the active monitoring IOM, while the other IOM becomes the standby. If the active monitoring IOM fails, the router manages the traffic switching to the standby IOM, and back to the active IOM, with minimum traffic lost.

The following figure shows redundant monitoring IOMs using a LAG.

Figure 61: Redundant monitoring IOMs using a LAG



Considerations for using redundant monitoring IOM solutions

Consider the following when implementing a redundant monitoring IOM solution:

- This solution is only supported in a single area; stitching of the FEC is not possible with this solution.
- In bandwidth monitoring with PXC ports, the arriving PMSI traffic must always be forwarded to the monitoring IOMs. If the monitoring IOM or IOMs reset or fail for any reason, the traffic and streams are interrupted.
- If no PXC is configured for bandwidth monitoring, the system uses legacy bandwidth monitoring on a single IOM. When a PXC with FPE is configured for bandwidth monitoring, the system only performs monitoring using the monitoring IOM, and the PXC configuration becomes mandatory for the system and UMH redundancy, whether the UMH redundancy is done on single ingress IOM or a double IOM.
- Any LDP policy assigned to the PXC interface is ignored during processing.
- MoFRR and ASBR-MoFRR are not supported for this solution.
- If the PXC interface is administratively disabled, the internal label assigned to the PXC interface is not deprogrammed. The PXC interface consumes the ILM label and maintains this label even if the PXC interface administrative status is down.
- Before configuring UMH rate monitoring, you must configure a provider tunnel as MLDP, under either the inclusive or selective PMSI context. The UMH rate monitoring configuration is blocked if the provider tunnel is not configured.

- The service must be shutdown before changing the PXC with FPE configuration.
- If an FPE is configured in the **configure service vprn mvpn provider-tunnel umh-rate-monitoring** context, you must remove the FPE from this context before deleting it from the **configure fwd-path-ext** context.
- To make any changes to the FPE (add, modify, or delete) in the VPRN service context, you must first administratively disable the PIM in the VPRN service context. This forces the PIM to withdraw the tunnels from MTTM or LDP, and add them back when it is administratively enabled again.

3.4.19.2 Configuring bandwidth monitoring without redundant IOMs

Prerequisites

- Before configuring UMH rate monitoring, you must configure a provider tunnel as MLDP, under either the inclusive or selective PMSI context. The UMH rate monitoring configuration is blocked if the provider tunnel is not configured.
- You must shutdown the associated VPRN service before making changes to the PXC with FPE configuration.

About this task

If you use the UMH redundancy solution with two IOMs using PXC ports, you must also use PXC if you want to implement the single IOM solution. This procedure describes the process for configuring bandwidth monitoring without redundant monitoring IOMs.

FPE configuration with port-based PXC, where the PXC consumes the port and that port is not usable any longer.

Procedure

Step 1. Configure a PXC port under a physical port.

If you are configuring PXC ports for a PXC LAG configuration, you must configure all the ports with the same breakout and port speed.

Example

MD-CLI

```
[ex:/configure port-xc]
A:admin@node-2#
  pxc 1 {
    admin-state enable
    port-id 1/1/c23/3
  }
  port 1/1/c23/3 {
    admin-state enable
  }
  port pxc-1.a {
    admin-state enable
  }
  port pxc-1.b {
    admin-state enable
  }
```

Example

classic CLI

```
A:node-2>config>port-xc# info
    pxc 1 create
        port 1/1/c23/3
        no shutdown
    exit
    port 1/1/c23/3
        no shutdown
    exit
    port pxc-1.a
        no shutdown
    exit
    port pxc-1.b
        no shutdown
    exit
```

- Step 2.** Assign the PXC port configured in step 1 to a multicast service using the following command, to create the Martian IP addresses required for PXC forwarding of multicast packets.

Example

MD-CLI

```
[ex:/configure fwd-path-ext]
A:admin@Dut-AC# info
    fpe 1 {
        path {
            pxc 1
        }
        application {
            mvpn-umh-rate-monitoring true
        }
    }
```

Example

classic CLI

```
A:node-2>config>fwd-path-ext# info
    fpe 1 create
        path pxc 1
        mvpn-umh-rate-monitoring
    exit
```

- Step 3.** Configure the FPE with a MAC chip-based PXC, to allow other applications to use the port.

Example

MD-CLI

```
*[pr:/configure]
A:admin@Dut-AC# info
    card 1 {
        mda 1 {
            admin-state enable
            xconnect {
                mac 1 {
                    loopback 1 {
                        bandwidth 100
                    }
                }
            }
        }
    }
```

```

    }
  }
}
port pxc-1.a {
  admin-state enable
}
port pxc-1.b {
  admin-state enable
}
port 1/1/ml/1 {
  port-xc {
    pxc 1 {
      admin-state enable
      port-id 1/1/ml/1
    }
  }
}
fwd-path-ext {
  fpe 1 {
    multi-path {
      path 1 {
        pxc 1
      }
    }
    application {
      mvpn-umh-rate-monitoring true
    }
  }
}
}

```

Example

classic CLI

```

A:node-2>config>card# info
mda 1
  xconnect
  mac 1 create
  loopback 1 create
  bandwidth 100
  exit
  exit
  exit
  no shutdown
  exit
exit
port 1/1/ml/1
  shutdown
  exit
port-xc
  pxc 1 create
  port 1/1/ml/1
  no shutdown
  exit
exit
port 1/1/ml/1
  no shutdown
  exit
port pxc-1.a
  no shutdown
  exit
port pxc-1.b
  no shutdown
  exit

```

```

fwd-path-ext
  fpe 1 create
    path pxc 1
    mvpn-umh-rate-monitoring
  exit
exit
exit

```

3.4.19.3 Configuring redundant monitoring using a LAG

Prerequisites

- Before configuring UMH rate monitoring, you must configure a provider tunnel as MLDP, under either the inclusive or selective PMSI context. The UMH rate monitoring configuration is blocked if the provider tunnel is not configured.
- You must shutdown the associated VPRN service before making changes to the PXC with FPE configuration.

About this task

Use two monitoring IOMs for monitoring redundancy. These monitoring IOMs can be the same as the ingress active or backup IOMs, or a separate IOM. Create a PXC port on each IOM and add it as a LAG member to a LAG bundle using the **per-link-hash** command. Use the **configure service vprn mvpn provider-tunnel inclusive umh-rate-monitoring fpe** command to configure the FPE type, and reference it using the **fwd-path-ext** command to create the interface on this LAG bundle with its own Martian IP address. When received at the ingress UMH IOMs, the PMSI traffic is forwarded to the active LAG IOM and, if there is a failure on the active LAG IOM, the PMSI traffic is forwarded to the standby IOM.



Note:

Bandwidth monitoring uses per-link hashing mode for the LAG. In this mode, the user cannot specifically configure the active and standby member. One of the IOMs in the redundant pair randomly becomes the active monitoring IOM, while the other IOM is the standby. If the active IOM fails, the router manages the traffic switching to the standby IOM, and then back to the active monitoring IOM, with minimum traffic lost.

For PXC LAG configuration, you must configure all the PXC ports with the same breakout and port speed.

Procedure

Step 1. Configure the PXC ports and FPE as described in [Configuring bandwidth monitoring without redundant IOMs](#).

For PXC LAG configuration, you must configure all the PXC ports with the same breakout and port speed.

Step 2. Configure the LAG.

Example

```

[ex:/configure lag "lag-10"]
A:admin@Dut-AC# info
  admin-state enable
  encap-type dot1q
  mode hybrid
  access {

```

```

        adapt-qos {
            mode link
        }
    }
    per-link-hash {
    }
    port pxc-1.a {
    }
    port pxc-2.a {
    }

[ex:/configure lag "lag-20"]
A:admin@Dut-AC# info
admin-state enable
encap-type dot1q
mode hybrid
access {
    adapt-qos {
        mode link
    }
}
per-link-hash {
}
port pxc-1.b {
}
port pxc-2.b {
}

```

Example classic CLI

```

A:node-2>config>lag# info
lag 10
mode hybrid
encap-type dot1q
access
    adapt-qos link
exit
port pxc-1.a
port pxc-2.a
per-link-hash
no shutdown
exit
lag 20
mode hybrid
encap-type dot1q
access
    adapt-qos link
exit
port pxc-1.b
port pxc-2.b
per-link-hash
no shutdown
exit

```

Step 3. Create an FPE to connect the PXC ports.

The PXC .a and .b subports must be added to the two LAGs being configured as **xc-a** and **xc-b** under the FPE. To see which actual port in the FPE is currently being used for rate monitoring, check the per-link-hash interface associations of the FPE's xc-a LAG using the **show lag 10 associations per-link-hash interface** command. See [Viewing the port used for monitoring](#) for more information.

Example

```
[ex:/configure port-xc]
A:admin@node-2#
  fwd-path-ext {
    fpe 1 {
      path {
        xc-lag-a "lag-10"
        xc-lag-b "lag-20"
      }
      application {
        mvpn-umh-rate-monitoring true
      }
    }
  }
}
```

Example**classic CLI**

```
A:node-2>config>port-xc# info
  fwd-path-ext
    fpe 1 create
      path xc-a lag-10 xc-b lag-20
      mvpn-umh-rate-monitoring
    exit
  exit
exit
```

3.4.19.4 Viewing the port used for monitoring**About this task**

To identify the port in the FPE that is currently being used for bandwidth monitoring, review the per-link hashing interface associations of the FPE's **xc-a** configured LAG.

Procedure

- Step 1.** Use the following command to view information about the specified LAG associations for a per-link hashing interface association of the FPE **xc-a** configured LAG.

```
show lag lag-id associations per-link-hash interface
```

Example

```
show lag 10 associations per-link-hash interface
```

- Step 2.** Review the output to determine the PXC IOM on which the monitoring is taking place. In the following example, the bandwidth monitoring takes place on the PXC 2 IOM.

Example

```
=====
Interface Associations
=====
Interface                               Active Link
-----
```

```

_tmnx_fpe_1.a          pxc-2.a
=====
Number of interface associations: 1
=====

```

3.5 FIB prioritization

The RIB processing of specific routes can be prioritized using the **rib-priority** command. This command allows specific routes to be prioritized through the protocol processing so that updates are propagated to the FIB as quickly as possible.

The **rib-priority** command can be configured within the VPRN instance of the OSPF or IS-IS routing protocols. For OSPF, a prefix list can be specified that identifies which route prefixes should be considered high priority.

Use the following command to configure all routes learned through the specified interface as high priority.

```

configure service vprn ospf area interface rib-priority high
configure service vprn ospf3 area interface rib-priority high

```

For the IS-IS routing protocol, RIB prioritization can be either specified through a prefix-list or an IS-IS tag value. If a prefix list is specified then route prefixes matching any of the prefix list criteria are considered high priority. If instead, an IS-IS tag value is specified, then any IS-IS route with that tag value is considered high priority.

The routes that have been designated as high priority are the first routes processed and then passed to the FIB update process so that the forwarding engine can be updated. All known high priority routes should be processed before the routing protocol moves on to other standard-priority routes. This feature has the most impact when there are a large number of routes being learned through routing protocols.

3.6 Configuring a VPRN service using CLI

This section provides information to configure VPRN services using the CLI.

3.6.1 Basic configuration

The following command options require specific input (there are no defaults) to configure a basic VPRN service:

- customer ID (see the *7705 SAR Gen 2 Services Overview Guide*)
- interface command options
- spoke SDP command options

The following example displays the configuration of a VPRN service.

Example: MD-CLI

```

[ex:/configure service vprn "32"]
A:admin@node-2# info
    admin-state enable
    customer "1"

```

```

autonomous-system 10000
ecmp 8
pim {
    apply-to all
    rp {
        ipv4 {
            bsr-candidate {
                admin-state disable
            }
        }
        ipv6 {
            bsr-candidate {
                admin-state disable
            }
        }
    }
}
bgp-ipvpn {
    mpls {
        admin-state enable
        route-distinguisher "10001:1"
        vrf-target {
            community "target:10001:1"
        }
        vrf-import {
            policy ["vrfImpPolCust1"]
        }
        vrf-export {
            policy ["vrfExpPolCust1"]
        }
        auto-bind-tunnel {
            resolution filter
            resolution-filter {
                ldp true
            }
        }
    }
}
bgp {
    router-id 10.0.0.1
    ebgp-default-reject-policy {
        import false
        export false
    }
    group "to-cel" {
        peer-as 65101
        export {
            policy ["vprnBgpExpPolCust1"]
        }
    }
    neighbor "10.1.1.2" {
        group "to-cel"
    }
}
interface "to-cel" {
    ipv4 {
        primary {
            address 10.1.0.1
            prefix-length 24
        }
        neighbor-discovery {
            proxy-arp-policy ["test"]
        }
    }
    dhcp {

```

```

        admin-state enable
        description "DHCP test"
    }
    vrrp 1 {
    }
}
sap 1/1/9:2 {
    ingress {
        qos {
            sap-ingress {
                policy-name "100"
            }
        }
    }
    egress {
        qos {
            sap-egress {
                policy-name "1010"
            }
        }
        filter {
            ip "10"
        }
    }
}
}
static-routes {
    route 10.5.0.0/24 route-type unicast {
        next-hop "10.1.1.2" {
            admin-state enable
        }
    }
}
rip {
    admin-state enable
    export-policy ["vprnRipExpPolCust1" "vprnRipExpoPolCust1"]
    group "cel" {
        admin-state enable
        neighbor "to-cel" {
            admin-state enable
        }
    }
    group "cel" {
        neighbor "to-cel" {
        }
    }
}
}

```

Example: classic CLI

```

A:node-2>config>service>vprn# info
-----
    ecmp 8
    autonomous-system 10000
    interface "to-cel" create
        address 10.1.0.1/24
        dhcp
            no shutdown
            description "DHCP test"
        exit
    vrrp 1
    exit
    proxy-arp-policy "test"

```

```

        sap 1/1/9:2 create
            ingress
                qos 100
            exit
            egress
                qos 1010
                filter ip 10
            exit
        exit
    exit
static-route-entry 10.5.0.0/24
    next-hop 10.1.1.2
    no shutdown
    exit
exit
bgp-ipvpn
    mpls
        auto-bind-tunnel
        resolution-filter
        ldp
        exit
        resolution filter
    exit
    route-distinguisher 10001:1
    vrf-import "vrfImpPolCust1"
    vrf-export "vrfExpPolCust1"
    vrf-target target:10001:1
    no shutdown
    exit
exit
bgp
    router-id 10.0.0.1
    group "to-cel"
        export "vprnBgpExpPolCust1"
        peer-as 65101
        neighbor 10.1.1.2
    exit
    exit
    no shutdown
exit
pim
    apply-to all
    rp
        static
        exit
        bsr-candidate
        shutdown
        exit
        rp-candidate
        shutdown
        exit
    exit
    no shutdown
exit
rip
    export "vprnRipExpPolCust1"
    group "cel"
        neighbor "to-cel"
        no shutdown
        exit
        no shutdown
    exit
    no shutdown
exit

```

```
no shutdown
```

3.6.2 Common configuration tasks

Prerequisites

This section provides a brief overview of the tasks that must be performed to configure a VPRN service and provides the CLI commands.

Perform the following tasks to configure a VPRN service.

Procedure

- Step 1.** Associate a VPRN service with a customer ID.
- Step 2.** Optional: Define an autonomous system.
- Step 3.** Define a route distinguisher.
- Step 4.** Define VRF route-target associations or VRF import/export policies.
- Step 5.** Optional: Define PIM command options.
- Step 6.** Optional: Create a subscriber interface.
- Step 7.** Create an interface.
- Step 8.** Define SAP command options on the interface.
 - a. Select nodes and ports.
 - b. Optional: Select QoS policies other than the default (configured in **configure qos** context).
 - c. Optional: Select filter policies (configured in **configure filter** context).
 - d. Optional: Select accounting policy (configured in **configure log** context).
 - e. Optional: Configure DHCP features.
- Step 9.** Optional: Define BGP command options.

BGP must be enabled in the **configure router bgp** context.
- Step 10.** Optional: Define RIP command options.
- Step 11.** Optional: Spoke SDP command options.
- Step 12.** Optional: Create confederation autonomous systems within an AS.
- Step 13.** Enable the service.

3.6.3 Configuring VPRN components

This section provides VPRN configuration examples.

3.6.3.1 Creating a VPRN service

Use the commands in the following context to create a VPRN service.

```
configure service vprn
```

A route distinguisher must be defined and the VPRN service must be administratively up in order for VPRN to be operationally active.

The following example displays a VPRN service configuration.

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2# info
  vprn "test" {
    admin-state enable
    service-id 1
    customer "1"
    ...
    admin-state enable
    route-distinguisher "10001:0"
  }
```

Example: classic CLI

```
A:node-2>config>service# info
-----
...
  vprn 1 name test customer 1 create
    route-distinguisher 10001:0
    no shutdown
  exit
...
-----
```

3.6.3.2 Configuring a global VPRN service

The following example displays a VPRN service with configured command options.

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2# info
  vprn "test" {
    admin-state enable
    service-id 27
    customer "1"
    autonomous-system 10000
    bgp-ipvprn {
      mpls {
        admin-state enable
        route-distinguisher "10001:1"
        vrf-import {
          policy ["vrfImpPolCust1"]
        }
        vrf-export {
          policy ["vrfExpPolCust1"]
        }
      }
    }
    spoke-sdp 2:27 {
    }
```

Example: classic CLI

```
A:node-2>config>service# info
...
vprn 27 name "test" customer 1 create
    autonomous-system 10000
    bgp-ipvpn
    mpls
    route-distinguisher 10001:1
    vrf-import "vrfImpPolCust1"
    vrf-export "vrfExpPolCust1"
    no shutdown
    exit
exit
spoke-sdp 2:27 create
```

3.6.3.3 Configuring a VPRN log

The following example displays a VPRN log configuration.

Example: MD-CLI

```
[ex:/configure service vprn "101"]
A:admin@node-2# info
customer "1"
log {
    filter "1" {
        named-entry "1" {
            action forward
        }
    }
    log-id "1" {
        filter "1"
        source {
            main true
            change true
        }
        destination {
            syslog "1"
        }
    }
    log-id "32" {
        filter "1"
        source {
            main true
            change true
        }
        destination {
            snmp {
            }
        }
    }
}
snmp-trap-group "32" {
    trap-target "3" {
        address 3ffe::e01:403
        port 9000
        version snmpv2c
        notify-community "vprn1"
    }
}
syslog "1" {
```

```

        address 3ffe::e01:403
        log-prefix "vprn1"
    }
}
snmp {
    access true
    community "dMHKqSM+0Ki7WFsaGl3Fy9Sn6wDeooe9Ltjrwvc5lw== hash2" {
        access-permissions r
    }
    community "80Ixno7a0LReeFUINhWfGeYS0vjzfLCX167ZYtjQp2o= hash2" {
        access-permissions rw
        version v2c
    }
}
dhcp-server {
    dhcpv4 "vprn_1" {
        admin-state enable
        force-renews true
        pool-selection {
            use-pool-from-client {
            }
        }
    }
}
}

```

Example: classic CLI

```

A:node-2>config>service>vprn# info
-----
    dhcp
        local-dhcp-server "vprn_1" create
            use-pool-from-client
            force-renews
            no shutdown
        exit
    exit
    snmp
        community "YsMv96H2KZVKQeakNAq.38gvyr.MH9vA" hash2 r version both
        community "gkYL94l90FFgu91PiRNvn3Rnl0edkMU1" hash2 rw version v2c
        access
    exit
    log
        filter 1 name "1"
            default-action forward
            entry 1 name "1"
                action forward
            exit
        exit
        syslog 1 name "1"
            address 3ffe::e01:403
            log-prefix "vprn1"
        exit
        snmp-trap-group 32 name "32"
            snmp-trap-group 32 name "32"
            trap-target "3" address 3ffe::e01:403 port 9000 snmpv2c notify-
community "vprn1"
        exit
        log-id 1 name "1"
            filter 1
            from main change
            to syslog 1
        exit
        log-id 32 name "32"

```

```

        filter 1
        from main change
        to snmp
        no shutdown
    exit
exit
-----

```

3.6.3.3.1 Configuring a spoke SDP

Use the commands in the following context to create or enable a spoke SDP.

```
configure service vprn spoke-sdp
```

Use the commands in the following context to configure the spoke-SDP command options.

```
configure service vprn interface spoke-sdp
```

3.6.3.4 Configuring VPRN protocols - PIM

The following example displays a VPRN PIM configuration.

Example: MD-CLI

```

[ex:/configure service]
A:admin@node-2# info
vprn "101" {
    admin-state enable
    customer "1"
    pim {
        apply-to all
        interface "if1" {
        }
        interface "if2" {
        }
        rp {
            ipv4 {
                bsr-candidate {
                    admin-state disable
                }
            }
            ipv6 {
                bsr-candidate {
                    admin-state disable
                }
            }
        }
    }
}
bgp-ipvpn {
    mpls {
        admin-state enable
        route-distinguisher "1:11"
    }
}
interface "if1" {
    loopback true
    ipv4 {

```

```

        primary {
            address 10.13.14.15
            prefix-length 32
        }
    }
}
interface "if2" {
    ipv4 {
        primary {
            address 10.13.14.1
            prefix-length 24
        }
    }
    sap 1/1/9:0 {
    }
}
}

```

Example: classic CLI

```

A:node-2>config>service# info
-----
vprn 101 name "101" customer 1 create
    interface "if1" create
        address 10.13.14.15/32
        loopback
    exit
    interface "if2" create
        address 10.13.14.1/24
        sap 1/1/9:0 create
    exit
exit
bgp-ipvpn
    mpls
        route-distinguisher 1:11
        no shutdown
    exit
exit
pim
    interface "if1"
    exit
    interface "if2"
    exit
    rp
        static
        exit
        bsr-candidate
            shutdown
        exit
        rp-candidate
            shutdown
        exit
    exit
    no shutdown
exit
no shutdown
exit

```

3.6.3.4.1 Configuring router interfaces

For the MD-CLI command descriptions and information to configure router interfaces, see the *7705 SAR Gen 2 MD-CLI Command Reference Guide*.

For the classic CLI command descriptions and information to configure router interfaces, see *7705 SAR Gen 2 Classic CLI Command Reference Guide*.

The following example displays a router interface configuration.

Example: MD-CLI

```
[ex:/configure service vprn "32"]
A:admin@node-2# info
...
  interface "if1" {
    port 1/1/33
    ipv4 {
      primary {
        address 10.2.2.1
        prefix-length 24
      }
    }
  }
  interface "if2" {
    port 1/1/34
    ipv4 {
      primary {
        address 10.49.1.46
        prefix-length 24
      }
    }
  }
}
```

Example: classic CLI

```
A:node-2>config>router# info
#-----
echo "IP Configuration"
#-----
...
  interface "if1"
    address 10.2.2.1/24
    port 1/1/33
  exit
  interface "if2"
    address 10.49.1.46/24
    port 1/1/34
  exit
```

3.6.3.4.2 Configuring VPRN protocols - BGP

The autonomous system number and router ID configured in the VPRN context only applies to that particular service.

The minimum command options that should be configured for a VPRN BGP instance are:

- Specify an autonomous system number for the router. See [Configuring a global VPRN service](#).

- Specify a router ID. If a new or different router ID value is entered in the BGP context, then the new value takes precedence and overwrites the VPRN-level router ID. See [Configuring a global VPRN service](#).
- Specify a VPRN BGP peer group.
- Specify a VPRN BGP neighbor with which to peer.
- Specify a VPRN BGP peer-AS that is associated with the above peer.

VPRN BGP is administratively enabled upon creation. There are no default VPRN BGP groups or neighbors. Each VPRN BGP group and neighbor must be explicitly configured.

All command options configured for VPRN BGP are applied to the group and are inherited by each peer, but a group command option can be overridden on a specific basis. VPRN BGP command hierarchy consists of three levels:

- global

```
configure service vprn bgp
```

- group

```
configure service vprn bgp group
```

- neighbor

```
configure service vprn bgp neighbor
```

The local address must be explicitly configured if two systems have multiple BGP peer sessions between them for the session to be established.

For more information about the BGP protocol, see the *7705 SAR Gen 2 Router Configuration Guide*.

3.6.3.4.2.1 Configuring VPRN BGP groups and neighbors

A group is a collection of related VPRN BGP peers. The group name should be a descriptive name for the group. Follow your group, name, and ID naming conventions for consistency and to help when troubleshooting faults.

All command options configured for a peer group are applied to the group and are inherited by each peer (neighbor), but a group command option can be overridden on a specific neighbor-level basis.

After a group name is created and options are configured, neighbors can be added within the same autonomous system to create IBGP connections or neighbors in different autonomous systems to create EBGP peers. All command options configured for the peer group level are applied to each neighbor, but a group command option can be overridden on a specific neighbor basis.

3.6.3.4.2.2 Configuring route reflection

Route reflection can be implemented in autonomous systems with a large internal BGP mesh to reduce the number of IBGP sessions required. One or more routers can be selected to act as focal points, for internal BGP sessions. Several BGP-speaking routers can peer with a route reflector. A route reflector forms peer connections to other route reflectors. A router assumes the role as a route reflector by configuring the commands in the **cluster** context. No other command is required unless you want to disable reflection to specific peers.

If you configure the cluster command at the global level, then all subordinate groups and neighbors are members of the cluster. The route reflector cluster ID is expressed in dotted-decimal notation. The ID should be a significant topology-specific value. No other command is required unless you want to disable reflection to specific peers.

If a route reflector client is fully meshed, use the following command to stop the route reflector from reflecting redundant route updates to a client.

- **MD-CLI**

```
configure router bgp client-reflect false
```

- **classic CLI**

```
configure router bgp disable-client-reflect
```

3.6.3.4.2.3 Configuring BGP confederations

A VPRN can be configured to belong to a BGP confederation. BGP confederations are one technique for reducing the degree of IBGP meshing within an AS. When the confederation command is in the configuration of a VPRN the type of BGP session formed with a VPRN BGP neighbor is determined as follows:

- The session is of type IBGP if the peer AS is the same as the local AS.
- The session is of type confed-EBGP if the peer AS is different than the local AS AND the peer AS is listed as one of the members in the confederation command.
- The session is of type EBGP if the peer AS is different than the local AS AND the peer AS is not listed as one of the members in the confederation command.

3.6.3.4.2.4 VPRN BGP configuration

The following example displays a VPRN BGP configuration.

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2#
    sdp 2 {
        admin-state enable
        keep-alive {
            admin-state disable
        }
        far-end {
            ip-address 1.2.3.4
        }
    }
...
    vprn "27" {
        customer "1"
        autonomous-system 10000
        ecmp 8
        bgp-ipvpn {
            mpls {
                admin-state enable
                route-distinguisher "10001:1"
            }
        }
    }
```

```

        vrf-target {
            community "target:10001:1"
        }
        vrf-import {
            policy ["vrfImpPolCust1"]
        }
        vrf-export {
            policy ["vrfExpPolCust1"]
        }
        auto-bind-tunnel {
            resolution filter
            resolution-filter {
                ldp true
            }
        }
    }
}
bgp {
    router-id 10.0.0.1
    ebgp-default-reject-policy {
        import false
        export false
    }
    group "to-cel" {
        peer-as 65101
        export {
            policy ["vprnBgpExpPolCust1"]
        }
    }
    neighbor "10.1.1.2" {
        group "to-cel"
    }
}
interface "to-cel" {
    ipv4 {
        primary {
            address 10.1.0.1
            prefix-length 24
        }
    }
    sap 1/1/9:2 {
        ingress {
            qos {
                sap-ingress {
                    policy-name "100"
                }
                scheduler-policy {
                    policy-name "SLA2"
                }
            }
        }
        egress {
            qos {
                sap-egress {
                    policy-name "1010"
                }
                scheduler-policy {
                    policy-name "SLA1"
                }
            }
            filter {
                ip "6"
            }
        }
    }
}

```

```

    }
  }
  spoke-sdp 2:27 {
  }
  static-routes {
    route 10.5.0.0/24 route-type unicast {
      next-hop "10.1.1.2" {
      }
    }
  }
}

```

Example: classic CLI

```

A:node-2>config>service# info
-----
sdp 2 create
  far-end 1.2.3.4
  keep-alive
  shutdown
  exit
  no shutdown
exit
customer 1 name "1" create
  description "Default customer"
exit
ipipe 1 name "1" customer 1 create
  shutdown
exit
vprn 27 name "27" customer 1 create
  shutdown
  ecmp 8
  autonomous-system 10000
  interface "to-cel" create
    address 10.1.0.1/24
    sap 1/1/9:2 create
      ingress
        scheduler-policy "SLA2"
        qos 100
      exit
      egress
        scheduler-policy "SLA1"
        qos 1010
        filter ip 6
      exit
    exit
  exit
  static-route-entry 10.5.0.0/24
    next-hop 10.1.1.2
    shutdown
  exit
exit
bgp-ipvpn
  mpls
    auto-bind-tunnel
    resolution-filter
      ldp
    exit
    resolution filter
  exit
  route-distinguisher 10001:1
  vrf-import "vrfImpPolCust1"
  vrf-export "vrfExpPolCust1"

```

```

        vrf-target target:10001:1
        no shutdown
    exit
exit
bgp
    router-id 10.0.0.1
    group "to-cel"
        export "vprnBgpExpPolCust1"
        peer-as 65101
        neighbor 10.1.1.2
    exit
exit
    no shutdown
exit
spoke-sdp 2:27 create
exit
exit

```

3.6.3.4.3 Configuring VPRN protocols - RIP

PE routers attached to a specific VPN must learn the set of addresses for each site in that VPN. There are several ways for a PE router to obtain this information, one of which is the Routing Information Protocol (RIP). RIP sends routing update messages that include entry changes to the routing table, which is updated with the new information.

RIP can be used as a PE and CE distribution technique. PE and CE routers can be configured as RIP peers, and the CE router can transmit RIP updates to inform the PE router about the set of address prefixes that are reachable at the CE router site.



Note: When RIP is configured in the CE router, care must be taken to ensure that address prefixes from other sites (address prefixes learned by the CE router from the PE router) are never advertised to the PE router. Specifically, if a PE router receives a VPN-IPv4 route and, as a result, distributes an IPv4 route to a CE router, that route must not be distributed back from that CE site to a PE router (either the same router or different routers).

Use the commands in the following context to enable a VPRN RIP instance and enable the RIP protocol.

```
configure service vprn rip
```

VPRN RIP is administratively enabled upon creation. Configuring other RIP commands and command options is optional.



Caution: Careful planning is essential to implement commands that can affect the behavior of VPRN RIP global, group, and neighbor levels. Because the RIP commands are hierarchical, analyze the values that can disable features on a particular level.

The command options configured at the VPRN RIP global level are inherited by the group and neighbor levels. Most hierarchical VPRN RIP commands can be modified on different levels; the most specific value is used. That is, a VPRN RIP group-specific command takes precedence over a global VPRN RIP command. A neighbor-specific command takes precedence over a global VPRN RIP and group-specific command. For example, if you modify a VPRN RIP neighbor-level command default, the new value takes precedence over VPRN RIP group- and global-level settings. VPRN RIP groups and neighbors are not created by default and each must be explicitly configured.

The minimal command options that should be configured for a VPRN instance include the following:

- Specify a VPRN RIP peer group.
- Specify a VPRN RIP neighbor with which to peer.
- Specify a VPRN RIP peer-AS that is associated with the above peer.

The VPRN RIP command hierarchy consists of three levels:

- global
- group
- neighbor

Example: MD-CLI

```
[ex:/configure service vprn "1" rip]
A:admin@node-2# info
  group "RIP-ALU-A" {
    neighbor "to-ALU-4" {
    }
  }
}
```

Example: classic CLI

```
A:node-2>config>router>rip# info
-----
  group "RIP-ALU-A"
    neighbor "to-ALU-4"
    exit
  exit
-----
```

3.6.3.4.3.1 Configuring VPRN RIP

The following example displays a VPRN RIP configuration.

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2# info
...
  vprn "1" {
    admin-state enable
    customer "1"
    autonomous-system 10000
    ecmp 8
    bgp-ipvprn {
      mpls {
        admin-state enable
        route-distinguisher "10001:1"
        vrf-target {
          community "target:1001:1"
        }
        vrf-import {
          policy ["vrfImpPolCust1"]
        }
        vrf-export {
          policy ["vrfExpPolCust1"]
        }
        auto-bind-tunnel {
```

```

        resolution filter
        resolution-filter {
            ldp true
        }
    }
}
}
bgp {
    router-id 10.0.0.1
    ebgp-default-reject-policy {
        import false
        export false
    }
    group "to-cel" {
        peer-as 65101
        export {
            policy ["vprnBgpExpPolCust1"]
        }
    }
    neighbor "10.1.1.2" {
        group "to-cel"
    }
}
interface "to-cel" {
    ipv4 {
        primary {
            address 10.1.0.1
            prefix-length 24
        }
    }
    sap 1/1/10:1 {
        ingress {
            qos {
                sap-ingress {
                    policy-name "100"
                }
                scheduler-policy {
                    policy-name "SLA2"
                }
            }
        }
        egress {
            qos {
                sap-egress {
                    policy-name "1010"
                }
                scheduler-policy {
                    policy-name "SLA1"
                }
            }
            filter {
                ip "6"
            }
        }
    }
}
}
spoke-sdp 2:1 {
}
static-routes {
    route 10.5.0.0/24 route-type unicast {
        next-hop "10.1.1.2"
        admin-state enable
    }
}
}

```

```

    }
    rip {
        admin-state enable
        export-policy ["vprnRipExpPolCust1"]
        group "cel" {
            admin-state enable
            neighbor "to-cel" {
                admin-state enable
            }
        }
    }
}

```

Example: classic CLI

```

A:node-2>config>service# info
-----
...
    vprn 1 name "1" customer 1 create
        ecmp 8
        autonomous-system 10000
        interface "to-cel" create
            address 10.1.0.1/24
            sap 1/1/10:1 create
                ingress
                    scheduler-policy "SLA2"
                    qos 100
                exit
                egress
                    scheduler-policy "SLA1"
                    qos 1010
                    filter ip 6
                exit
            exit
        static-route-entry 10.5.0.0/24
            next-hop 10.1.1.2
            no shutdown
        exit
        bgp-ipvpn
            mpls
                auto-bind-tunnel
                resolution-filter
                ldp
            exit
            resolution filter
        exit
        route-distinguisher 10001:1
        vrf-import "vrfImpPolCust1"
        vrf-export "vrfExpPolCust1"
        vrf-target target:1001:1
        no shutdown
    exit
    exit
    bgp
        router-id 10.0.0.1
        group "to-cel"
            export "vprnBgpExpPolCust1"
            peer-as 65101
            neighbor 10.1.1.2
        exit
    exit

```

```
        no shutdown
    exit
    rip
        export "vprnRipExpPolCust1"
        group "cel"
            neighbor "to-cel"
            no shutdown
        exit
        no shutdown
    exit
    no shutdown
    exit
    spoke-sdp 2:1 create
    exit
    no shutdown
exit
...
-----
```

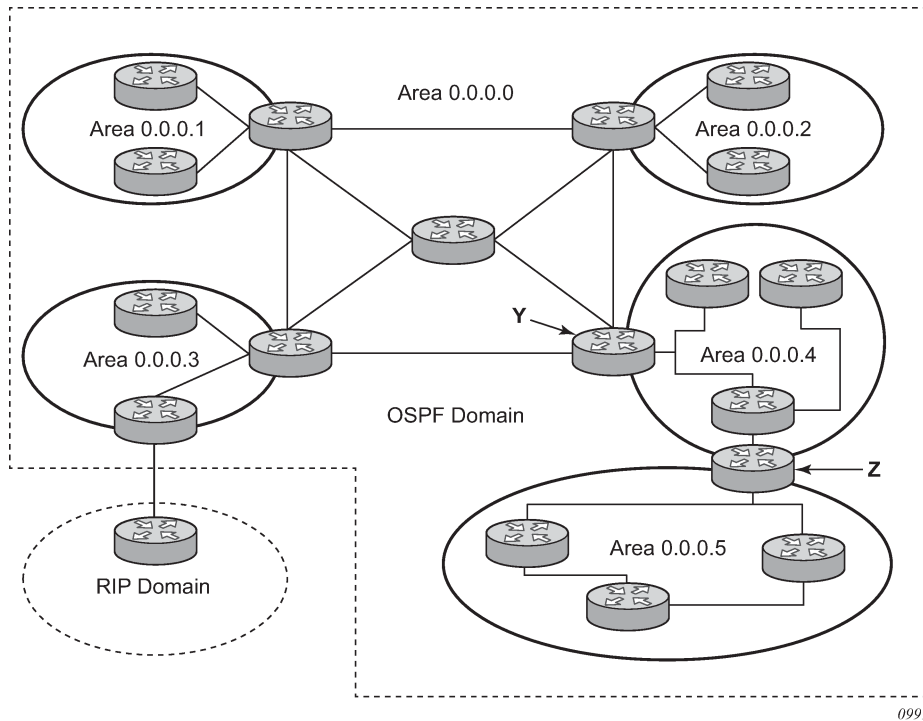
For more information about the RIP protocol, see the *7705 SAR Gen 2 Router Configuration Guide*.

3.6.3.4.4 Configuring VPRN protocols - OSPF

Each VPN routing instance is isolated from any other VPN routing instance, and from the routing used across the backbone. OSPF can be run with any VPRN, independently of the routing protocols used in other VPRNs, or in the backbone itself. For more information about the OSPF protocol, see the *7705 SAR Gen 2 Router Configuration Guide*.

The OSPF backbone area, area 0.0.0.0, must be contiguous and all other areas must be connected to the backbone area. The backbone distributes routing information between areas. If it is not practical to connect an area to the backbone (see Area 0.0.0.5 in [Figure 62: OSPF areas](#)), the area border routers (such as routers Y and Z) must be connected via a virtual link. The two area border routers form a point-to-point-like adjacency across the transit area (see Area 0.0.0.4). A virtual link can only be configured while in the area 0.0.0.0 context.

Figure 62: OSPF areas



3.6.3.4.4.1 Configuring VPRN OSPF

The following example displays the VPRN OSPF configuration shown in [Figure 62: OSPF areas](#).

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2# info
...
  vprn "1" {
    admin-state enable
    customer "1"
    interface "test" {
    }
    ospf 0 {
      admin-state enable
      area 0.0.0.0 {
        virtual-link 1.2.3.4 transit-area 1.2.3.4 {
          hello-interval 9
          dead-interval 40
        }
      }
      area 1.2.3.4 {
      }
    }
  }
}
```

Example: classic CLI

```

A:node-2>config>service# info
-----
...
    vprn 1 name "1" customer 1 create
        interface "test" create
        exit
        ospf
            no shutdown
            area 0.0.0.0
                virtual-link 1.2.3.4 transit-area 1.2.3.4
                hello-interval 9
                dead-interval 40
            exit
        exit
        area 1.2.3.4
        exit
    exit
    no shutdown
    exit
...
-----

```

For more information about the OSPF protocol, see the *7705 SAR Gen 2 Router Configuration Guide*.

3.6.3.4.5 Configuring a VPRN interface

Interface names associate an IP address to the interface, and then associate the IP interface with a physical port. The logical interface can associate attributes like an IP address, port, Link Aggregation Group (LAG) or the system.

There are no default interfaces.

You can configure a VPRN interface as a loopback interface by issuing the loopback command instead of the **sap** command. The loopback flag cannot be set on an interface where a SAP is already defined and a SAP cannot be defined on a loopback interface.

When using mtrace or mstat in a Layer 3 VPN context then the configuration for the VPRN should have a loopback address configured which has the same address as the core instance's system address (BGP next-hop).

The following example displays a VPRN interface configuration.

Example: MD-CLI

```

[ex:/configure service]
A:admin@node-2# info
...
    vprn "1" {
        admin-state enable
        customer "1"
        autonomous-system 10000
        ecmp 8
        bgp-ipvpn {
            mpls {
                admin-state enable
                route-distinguisher "10001:1"
                vrf-target {
                    community "target:1001:1"
                }
            }
        }
    }

```

```

    }
    vrf-import {
        policy ["vrfImpPolCust1"]
    }
    vrf-export {
        policy ["vrfExpPolCust1"]
    }
    auto-bind-tunnel {
        resolution filter
        resolution-filter {
            ldp true
        }
    }
}
}
interface "to-cel" {
    ipv4 {
        primary {
            address 10.1.0.1
            prefix-length 24
        }
    }
}
spoke-sdp 2:1 {
}
static-routes {
    route 10.5.0.0/24 route-type unicast {
        next-hop "10.1.1.2" {
            admin-state enable
        }
    }
}
}
...

```

Example: classic CLI

```

A:node-2>config>service# info
-----
...
    vprn 1 name "1" customer 1 create
        ecmp 8
        autonomous-system 10000
        interface "to-cel" create
            address 10.1.0.1/24
        exit
        static-route-entry 10.5.0.0/24
            next-hop 10.1.1.2
            no shutdown
        exit
        exit
        bgp-ipvpn
            mpls
                auto-bind-tunnel
                    resolution-filter
                        ldp
                exit
                resolution filter
            exit
            route-distinguisher 10001:1
            vrf-import "vrfImpPolCust1"
            vrf-export "vrfExpPolCust1"
            vrf-target target:1001:1

```

```

        no shutdown
    exit
exit
spoke-sdp 2:1 create
exit
no shutdown
exit
...
-----

```

3.6.3.4.6 Configuring a VPRN interface SAP

A SAP is a combination of a port and encapsulation command options that identifies the service access point on the interface and within the SR. Each SAP must be unique within a router. A SAP cannot be defined if the interface **loopback** command is enabled.

When configuring VPRN interface SAP command options, a default QoS policy is applied to each ingress and egress SAP. Additional QoS policies and scheduler policies must be configured in the **configure qos** context. Filter policies are configured in the **configure filter** context and must be explicitly applied to a SAP. There are no default filter policies.

The following example displays a VPRN interface SAP configuration.

Example: MD-CLI

```

[ex:/configure service]
A:admin@node-2# info
...
    vprn "1" {
        admin-state enable
        customer "1"
        autonomous-system 10000
        ecmp 8
        bgp-ipvprn {
            mpls {
                admin-state enable
                route-distinguisher "10001:1"
                vrf-target {
                    community "target:1001:1"
                }
            }
            vrf-import {
                policy ["vrfImpPolCust1"]
            }
            vrf-export {
                policy ["vrfExpPolCust1"]
            }
            auto-bind-tunnel {
                resolution filter
                resolution-filter {
                    ldp true
                }
            }
        }
    }
    interface "to-cel" {
        ipv4 {
            primary {
                address 10.1.0.1
                prefix-length 24
            }
        }
    }

```

```

    }
    sap 1/1/10:1 {
        ingress {
            qos {
                sap-ingress {
                    policy-name "100"
                }
                scheduler-policy {
                    policy-name "SLA2"
                }
            }
        }
        egress {
            qos {
                sap-egress {
                    policy-name "1010"
                }
                scheduler-policy {
                    policy-name "SLA1"
                }
            }
            filter {
                ip "6"
            }
        }
    }
}
spoke-sdp 2:1 {
}
static-routes {
    route 10.5.0.0/24 route-type unicast {
        next-hop "10.1.1.2" {
            admin-state enable
        }
    }
}
}

```

Example: classic CLI

```

A:node-2>config>service# info
-----
...
    vprn 1 name "1" customer 1 create
        ecmp 8
        autonomous-system 10000
        interface "to-cel" create
            address 10.1.0.1/24
            sap 1/1/10:1 create
                ingress
                    scheduler-policy "SLA2"
                    qos 100
                exit
                egress
                    scheduler-policy "SLA1"
                    qos 1010
                    filter ip 6
                exit
            exit
        exit
    static-route-entry 10.5.0.0/24
        next-hop 10.1.1.2
        no shutdown
    exit

```

```

    exit
    bgp-ipvpn
    mpls
        auto-bind-tunnel
        resolution-filter
        ldp
        exit
        resolution filter
    exit
    route-distinguisher 10001:1
    vrf-import "vrfImpPolCust1"
    vrf-export "vrfExpPolCust1"
    vrf-target target:1001:1
    no shutdown
    exit
    exit
    spoke-sdp 2:1 create
    exit
    no shutdown
    exit

```

3.7 Service management tasks

This section describes the VPRN service management tasks.

3.7.1 Modifying a VPRN service

The following example shows a VPRN service configuration that is used in the following sections to show how to modify the configuration.

Example: MD-CLI

```

[ex:/configure service]
A:admin@node-2# info
...
    vprn "1" {
        admin-state enable
        customer "1"
        autonomous-system 10000
        ecmp 8
        bgp-ipvpn {
            mpls {
                admin-state enable
                vrf-import {
                    policy ["vrfImpPolCust1"]
                }
                vrf-export {
                    policy ["vrfExpPolCust1"]
                }
            }
        }
    }
    bgp {
        router-id 10.0.0.1
        ebgp-default-reject-policy {
            import false
            export false
        }
        group "to-cel" {

```

```

        admin-state enable
        peer-as 65101
        export {
            policy ["vprnBgpExpPolCust1"]
        }
    }
    neighbor "10.1.1.2" {
        group "to-cel"
    }
}
maximum-ipv4-routes {
    value 2000
}
interface "to-cel" {
    ipv4 {
        primary {
            address 10.1.1.1
            prefix-length 24
        }
    }
    sap 1/1/10:1 {
    }
}
spoke-sdp 2:1 {
}
static-routes {
    route 10.5.0.0/24 route-type unicast {
        next-hop "10.1.1.2" {
            admin-state enable
        }
    }
}
}

```

Example: classic CLI

```

A:node-2>config>service# info
-----
...
vprn 1 name "1" customer 1 create
no shutdown
ecmp 8
maximum-routes 2000
autonomous-system 10000
interface "to-cel" create
    address 10.1.1.1/24
    sap 1/1/10:1 create
    exit
exit
static-route-entry 10.5.0.0/24
    next-hop 10.1.1.2
    no shutdown
    exit
exit
bgp-ipvpn
    mpls
        no shutdown
        vrf-import "vrfImpPolCust1"
        vrf-export "vrfExpPolCust1"
    exit
exit
bgp
    router-id 10.0.0.1
    group "to-cel"

```

```

        export "vprnBgpExpPolCust1"
        peer-as 65101
        neighbor 10.1.1.2
        exit
    exit
    no shutdown
    exit
    spoke-sdp 2:1 create
    exit
exit

```

3.7.2 Deleting a VPRN service

The following example displays the deletion of a VPRN service.

Example: MD-CLI

```

[ex:/configure service vprn "1"]
A:admin@node-2# exit

[ex:/configure service]
A:admin@node-2# delete vprn 1

```

Example: classic CLI

In the classic CLI, a VPRN service cannot be deleted until SAPs and interfaces are administratively disabled and deleted. If protocols or a spoke-SDP, or both are defined, they must be shut down and removed from the configuration as well.

```

*A:node-2>config>service# vprn 1
*A:node-2>config>service>vprn# interface "to-cel"
*A:node-2>config>service>vprn>if# sap 1/1/10:1
*A:node-2>config>service>vprn>if>sap# shutdown
*A:node-2>config>service>vprn>if>sap# exit
*A:node-2>config>service>vprn>if# no sap 1/1/10:1
*A:node-2>config>service>vprn>if# shutdown
*A:node-2>config>service>vprn>if# exit
*A:node-2>config>service>vprn# no interface "to-cel"
*A:node-2>config>service>vprn# bgp
*A:node-2>config>service>vprn>bgp# shutdown
*A:node-2>config>service>vprn>bgp# exit
*A:node-2>config>service>vprn# no bgp
*A:node-2>config>service>vprn# rip
*A:node-2>config>service>vprn>rip$ shutdown
*A:node-2>config>service>vprn>rip$ exit
*A:node-2>config>service>vprn# no rip
*A:node-2>config>service>vprn# no spoke-sdp 2
*A:node-2>config>service>vprn# no ecmp
*A:node-2>config>service>vprn# static-route-entry 10.5.0.0/24
*A:node-2>config>service>vprn>static-route-entry# no next-hop 10.1.1.2
*A:node-2>config>service>vprn>static-route-entry# shutdown
*A:node-2>config>service>vprn>static-route-entry# exit
*A:node-2>config>service>vprn# no static-route-entry 10.5.0.0/24
*A:node-2>config>service>vprn# exit
*A:node-2>config>service# no vprn 1

```

3.7.3 Disabling a VPRN service

A VPRN service can be administratively disabled without deleting any service command options. The following example displays the disabling of a VPRN service.

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2# vprn 1

[ex:/configure service vprn "1"]
A:admin@node-2# admin-state disable

[ex:/configure service vprn "1"]
A:admin@node-2# info
    admin-state disable
    customer "1"
...
```

Example: classic CLI

```
A:node-2>config>service# vprn 1
A:node-2>config>service>vprn# shutdown
A:node-2>config>service>vprn# info
-----
    shutdown
    ecmp 8
...
```

3.7.4 Re-enabling a VPRN service

The following example displays the re-enabling of a VPRN service that had been administratively disabled.

Example: MD-CLI

```
[ex:/configure service]
A:admin@node-2# vprn 1

[ex:/configure service vprn "1"]
A:admin@node-2# admin-state enable

[ex:/configure service vprn "1"]
A:admin@node-2# info
    admin-state enable
    customer "1"
...
```

Example: classic CLI

```
A:node-2>config>service# vprn 1
A:node-2>config>service>vprn# no shutdown
A:node-2>config>service>vprn# info
-----
    no shutdown
    ecmp 8
```

...

4 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.

4.1 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)*

4.2 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*

4.3 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*

4.4 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*

4.5 Ethernet

IEEE 802.3x, *Ethernet Flow Control*

4.6 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)*

4.7 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*

4.8 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*

4.9 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*

4.10 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*

4.11 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

4.12 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*

4.13 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*

4.14 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*

4.15 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*

4.16 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*

4.17 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*

4.18 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*

4.19 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*

4.20 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*

4.21 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)*

4.22 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*

4.23 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*

4.24 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*

4.25 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*

4.26 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*

4.27 Storage sanitization

IEEE 2883-2022, *IEEE Standard for Sanitizing Storage* – clear sanitization method for CF, MMC, SSD, SD, USB

4.28 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*

4.29 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

4.30 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)*

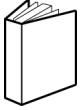
4.31 Yet Another Next Generation (YANG)

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](#)