



Configuring BGP Services for VOSS

Release 8.0.50 (VOSS for XA1400 Series)
9035651-02 Rev AA
August 2019

© 2017-2019, Extreme Networks, Inc.
All Rights Reserved.

Legal Notice

Extreme Networks, Inc. reserves the right to make changes in specifications and other information contained in this document and its website without prior notice. The reader should in all cases consult representatives of Extreme Networks to determine whether any such changes have been made.

The hardware, firmware, software or any specifications described or referred to in this document are subject to change without notice.

Trademarks

Extreme Networks and the Extreme Networks logo are trademarks or registered trademarks of Extreme Networks, Inc. in the United States and/or other countries.

All other names (including any product names) mentioned in this document are the property of their respective owners and may be trademarks or registered trademarks of their respective companies/owners.

For additional information on Extreme Networks trademarks, please see: www.extremenetworks.com/company/legal/trademarks

Open Source Declarations

Some software files have been licensed under certain open source or third-party licenses. End-user license agreements and open source declarations can be found at: www.extremenetworks.com/support/policies/software-licensing

Contents

Chapter 1: About this Document	6
Purpose	6
Conventions	6
Text Conventions	7
Documentation and Training	9
Getting Help	9
Providing Feedback to Us	10
Chapter 2: New in this Document	11
Notice about Feature Support	11
Chapter 3: BGP fundamentals	12
Autonomous systems	12
BGP 4 Byte AS Support	16
Routing information consolidation	19
BGP communities	27
BGP path attributes	28
BGP Route Selection	28
BGP and dampened routes	30
BGP Updates	31
Equal Cost Multipath	34
MD5 message authentication	35
BGP and route redistribution	36
BGP+	37
ECMP with BGP+	38
BGPv6	39
Circuitless IP	41
BGP Configuration Considerations and Limitations	42
Chapter 4: BGP configuration using CLI	48
Configuring BGP Globally	48
Configuring 4-byte AS numbers	53
Configuring Aggregate Routes	55
Configuring allowed networks	57
Configuring BGP peers or peer groups	58
Configuring a BGP peer or peer group password	63
Configuring redistribution to BGP	64
Configuring redistribution to BGP+ for VRF 0	69
Configuring AS path lists	72
Configuring community lists	73
Configuring extended community lists	74
Chapter 5: BGP verification using CLI	76

Viewing BGP aggregate information.....	76
Viewing IPv6 BGP+ aggregate information.....	76
Viewing CIDR routes.....	77
Viewing BGP configuration.....	78
Viewing BGP confederation.....	79
Viewing flap-dampened routes.....	80
Viewing global flap-dampening configurations.....	81
Viewing imported routes.....	82
Viewing BGPv6 imported routes.....	84
Viewing BGP neighbors information.....	84
Viewing BGPv6 neighbors information.....	86
Viewing BGP network configurations.....	88
Viewing IPv6 BGP+ network configurations.....	88
Viewing BGP peer group information.....	89
Viewing BGP redistributed routes.....	90
Viewing BGPv6 redistributed routes.....	91
Viewing a summary of BGP configurations.....	91
Viewing a summary of BGPv6 configurations.....	93
Viewing BGP routes.....	94
Viewing BGPv6 routes.....	95
Chapter 6: BGP configuration using EDM.....	97
Configuring BGP globally.....	97
Configuring 4-byte AS numbers.....	100
Configuring aggregate routes.....	102
Configuring aggregate IPv6 routes.....	103
Configuring allowed networks.....	104
Configuring allowed IPv6 networks.....	105
Configuring BGP peers.....	106
Configuring BGPv6 peers.....	109
Configuring peer groups.....	113
Viewing IPv6 Community Attributes.....	115
Displaying dampened routes information.....	116
Configuring redistribution to BGP.....	117
Configuring redistribution to BGPv6.....	118
Viewing BGP+ or BGPv6 route summary information.....	119
Viewing BGP route summary.....	120
Configuring an AS path list.....	121
Configuring a community access list.....	122
Chapter 7: BGP Configuration Examples.....	123
IPv6 Tunnel Configurations for BGP+.....	123
eBGP+ peership between two switches with IPv6 Tunneling.....	125
iBGP+ peership on CLIP between two switches with IPv6 Tunneling.....	126
Native IPv6 eBGP peership between two switches on VRF.....	129

Glossary..... 134

Chapter 1: About this Document

This section discusses the purpose of this document, the conventions used, ways to provide feedback, additional help, and information regarding other Extreme Networks publications.

Purpose

This document provides information on features in VSP Operating System Software (VOSS). VOSS runs on the following product families:

- Extreme Networks VSP 4000 Series
- Extreme Networks VSP 7200 Series
- Extreme Networks VSP 7400 Series
- Extreme Networks VSP 8000 Series (includes VSP 8200 Series and VSP 8400 Series)
- Extreme Networks VSP 8600 Series
- Extreme Networks XA1400 Series

This document provides conceptual information and procedures that you can use to configure Border Gateway Protocol (BGP) services. The following operations are supported by BGP:

- IPv4
- 4-byte AS
- Peer groups
- Redistribution

Examples and network illustrations in this document may illustrate only one of the supported platforms. Unless otherwise noted, the concept illustrated applies to all supported platforms.

Conventions

This section discusses the conventions used in this guide.

Text Conventions

The following tables list text conventions that can be used throughout this document.

Table 1: Notice Icons

Icon	Alerts you to...
 Important:	A situation that can cause serious inconvenience.
 Note:	Important features or instructions.
 Tip:	Helpful tips and notices for using the product.
 Danger:	Situations that will result in severe bodily injury; up to and including death.
 Warning:	Risk of severe personal injury or critical loss of data.
 Caution:	Risk of personal injury, system damage, or loss of data.

Table 2: Text Conventions

Convention	Description
Angle brackets (< >)	Angle brackets (< >) indicate that you choose the text to enter based on the description inside the brackets. Do not type the brackets when you enter the command. If the command syntax is <code>cfm maintenance-domain maintenance-level <0-7></code>, you can enter <code>cfm maintenance-domain maintenance-level 4</code>.
Bold text	Bold text indicates the GUI object name you must act upon. Examples: - Click OK. - On the Tools menu, choose Options.
Braces ({ })	Braces ({ }) indicate required elements in syntax descriptions. Do not type the braces when you enter the command. For example, if the command syntax is <code>ip address {A.B.C.D}</code>, you must enter the IP address in dotted, decimal notation.

Table continues...

Convention	Description
Brackets ([])	Brackets ([]) indicate optional elements in syntax descriptions. Do not type the brackets when you enter the command. For example, if the command syntax is <code>show clock [detail]</code>, you can enter either <code>show clock</code> or <code>show clock detail</code>.
Ellipses (...)	An ellipsis (...) indicates that you repeat the last element of the command as needed. For example, if the command syntax is <code>ethernet/2/1 [<parameter> <value>] ...</code>, you enter <code>ethernet/2/1</code> and as many parameter-value pairs as you need.
<i>Italic Text</i>	Italics emphasize a point or denote new terms at the place where they are defined in the text. Italics are also used when referring to publication titles that are not active links.
Plain Courier Text	Plain Courier text indicates command names, options, and text that you must enter. Plain Courier text also indicates command syntax and system output, for example, prompts and system messages. Examples: <code>show ip route</code> <code>Error: Invalid command syntax</code> <code>[Failed][2013-03-22 13:37:03.303 -04:00]</code>
Separator (>)	A greater than sign (>) shows separation in menu paths. For example, in the Navigation tree, expand the Configuration > Edit folders.
Vertical Line ()	A vertical line () separates choices for command keywords and arguments. Enter only one choice. Do not type the vertical line when you enter the command. For example, if the command syntax is <code>access-policy by-mac action { allow deny }</code>, you enter either <code>access-policy by-mac action allow</code> or <code>access-policy by-mac action deny</code>, but not both.

Documentation and Training

To find Extreme Networks product guides, visit our documentation pages at:

Current Product Documentation	www.extremenetworks.com/documentation/
Archived Documentation (for earlier versions and legacy products)	www.extremenetworks.com/support/documentation-archives/
Release Notes	www.extremenetworks.com/support/release-notes
Hardware/Software Compatibility Matrices	https://www.extremenetworks.com/support/compatibility-matrices/
White papers, data sheets, case studies, and other product resources	https://www.extremenetworks.com/resources/

Training

Extreme Networks offers product training courses, both online and in person, as well as specialized certifications. For more information, visit www.extremenetworks.com/education/.

Getting Help

If you require assistance, contact Extreme Networks using one of the following methods:

<u>Extreme Portal</u>	Search the GTAC (Global Technical Assistance Center) knowledge base, manage support cases and service contracts, download software, and obtain product licensing, training, and certifications.
<u>The Hub</u>	A forum for Extreme Networks customers to connect with one another, answer questions, and share ideas and feedback. This community is monitored by Extreme Networks employees, but is not intended to replace specific guidance from GTAC.
<u>Call GTAC</u>	For immediate support: 1-800-998-2408 (toll-free in U.S. and Canada) or +1 408-579-2826. For the support phone number in your country, visit: www.extremenetworks.com/support/contact

Before contacting Extreme Networks for technical support, have the following information ready:

- Your Extreme Networks service contract number and/or serial numbers for all involved Extreme Networks products
- A description of the failure
- A description of any action(s) already taken to resolve the problem
- A description of your network environment (such as layout, cable type, other relevant environmental information)

- Network load at the time of trouble (if known)
- The device history (for example, if you have returned the device before, or if this is a recurring problem)
- Any related RMA (Return Material Authorization) numbers

Subscribing to Service Notifications

You can subscribe to email notifications for product and software release announcements, Vulnerability Notices, and Service Notifications.

1. Go to www.extremenetworks.com/support/service-notification-form.
2. Complete the form with your information (all fields are required).
3. Select the products for which you would like to receive notifications.

 Note:

You can modify your product selections or unsubscribe at any time.

4. Click **Submit**.

Providing Feedback to Us

Quality is our first concern at Extreme Networks, and we have made every effort to ensure the accuracy and completeness of this document. We are always striving to improve our documentation and help you work better, so we want to hear from you! We welcome all feedback but especially want to know about:

- Content errors or confusing or conflicting information.
- Ideas for improvements to our documentation so you can find the information you need faster.
- Broken links or usability issues.

If you would like to provide feedback to the Extreme Networks Information Development team, you can do so in two ways:

- Use our short online feedback form at <https://www.extremenetworks.com/documentation-feedback/>.
- Email us at documentation@extremenetworks.com.

Please provide the publication title, part number, and as much detail as possible, including the topic heading and page number if applicable, as well as your suggestions for improvement.

Chapter 2: New in this Document

There are no feature changes in this document.

Notice about Feature Support

This document includes content for multiple hardware platforms across different software releases. As a result, the content can include features not supported by your hardware in the current software release.

If a documented command, parameter, tab, or field does not appear on your hardware, it is not supported.

For information about feature support, see [Release Notes for XA1400 Series](#).

For information about physical hardware restrictions, see your hardware documentation.

Chapter 3: BGP fundamentals

Border Gateway Protocol (BGP) is an inter-domain routing protocol that provides loop-free routing between autonomous systems (AS) or within an AS. This section describes the major BGP features.

Autonomous systems

An Autonomous system (AS) is a group of routers and hosts run by a single technical administrator that has a single, clearly defined routing policy. Each AS uses a unique AS number assigned by the appropriate Internet Registry entity. LANs and WANs that interconnect by IP routers form a group of networks called an internetwork. For administrative purposes, internetworks divide into boundaries known as autonomous systems.

The following figure shows a sample internetwork segmented into three autonomous systems.

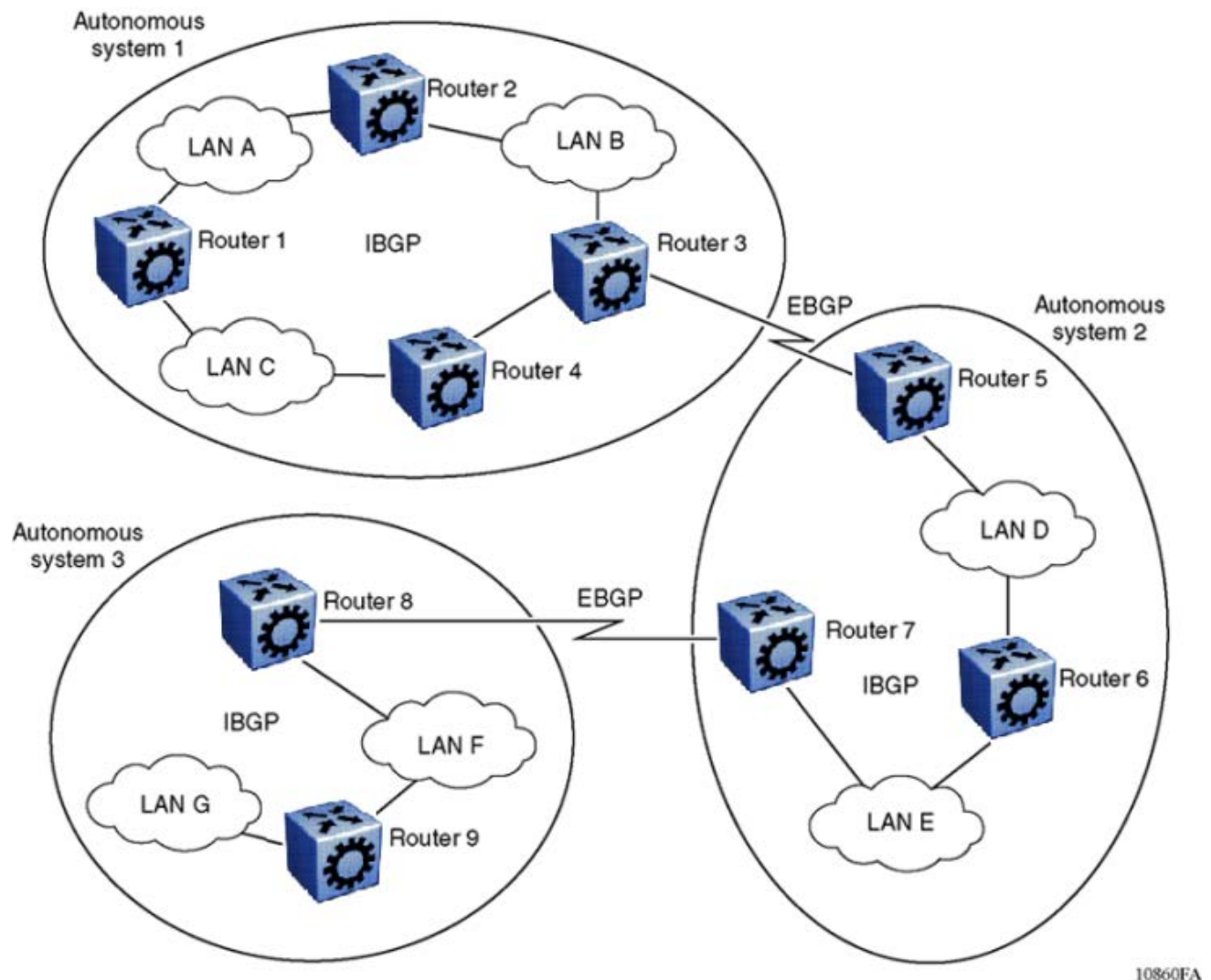


Figure 1: Internetwork segmented into three autonomous systems

BGP exchanges information between autonomous systems as well as between routers within the same AS. As shown in the preceding figure, routers that are members of the same AS and exchange BGP updates run internal BGP (iBGP), and routers that are members of different autonomous systems and exchange BGP updates run external BGP (eBGP).

Internal and external BGP routing

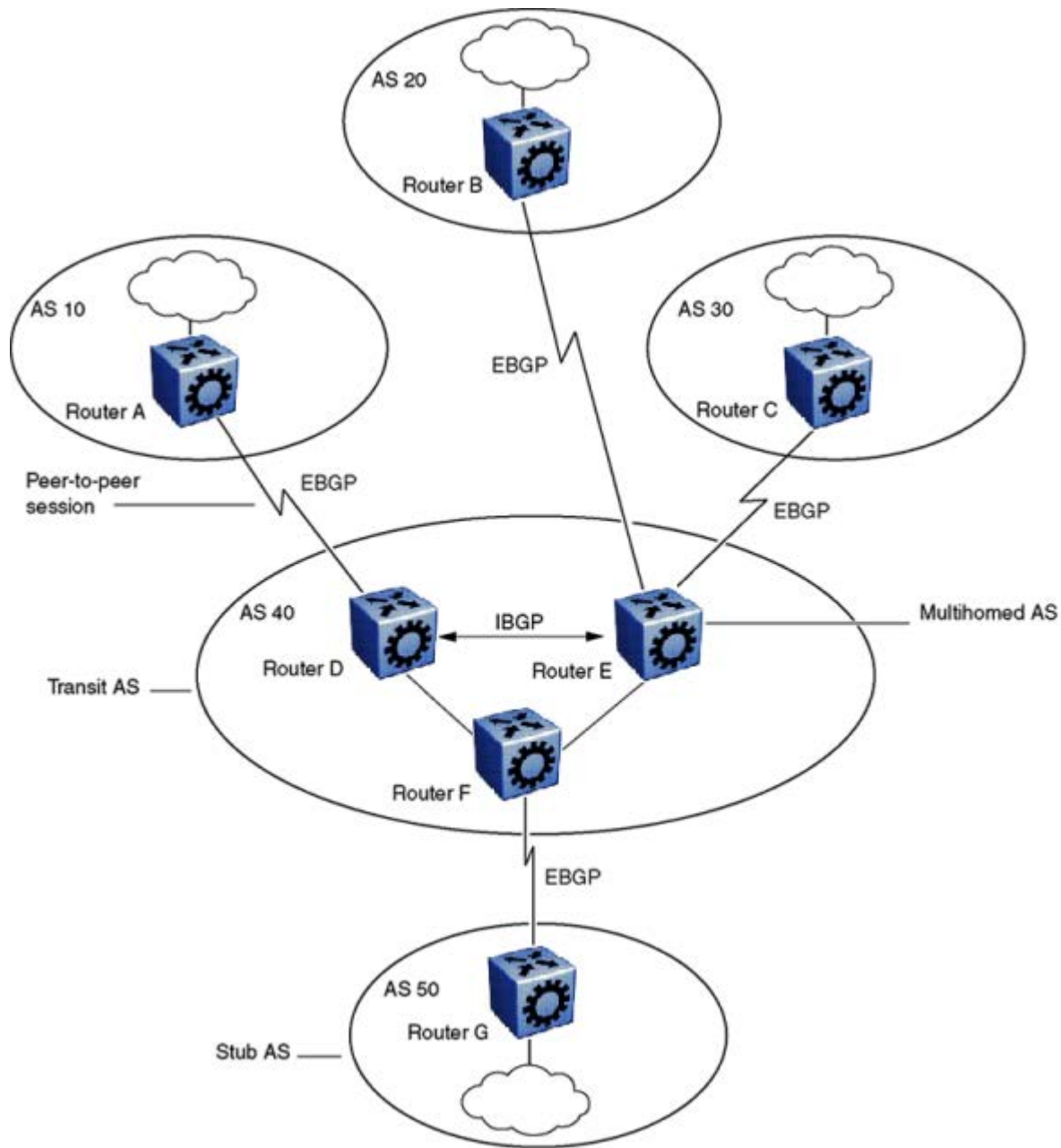
The switch supports both iBGP intra-AS routing and eBGP external-AS routing. With iBGP, each router within an AS runs an interior gateway protocol (IGP), such as Routing Information Protocol (RIP) or Open Shortest Path First (OSPF). The iBGP information, along with the IGP route to the originating BGP border router, determines the next hop to use to exchange information with an external AS. Each router uses iBGP exclusively to determine reachability to external autonomous systems. After a router receives an iBGP update destined for an external AS, it passes the update to IP for inclusion in the routing table only if a viable IGP route to the correct border gateway is available.

BGP speakers in different autonomous systems use eBGP to communicate routing information.

BGP speaker

BGP routers employ an entity within the router, referred to as a BGP speaker, which transmits and receives BGP messages and acts upon them. BGP speakers establish a peer-to-peer session with other BGP speakers to communicate.

All BGP speakers within an AS must be fully meshed. The following figure shows a BGP network with fully-meshed BGP speakers.



10861FA

Figure 2: BGP networks

Transit AS

An AS with more than one BGP speaker can use iBGP to provide a transit service for networks located outside the AS. An AS that provides this service is a transit AS. As shown in the preceding figure, [BGP networks](#) on page 14, AS 40 is the transit AS. AS 40 provides information about the internal networks, as well as transit networks, to the remaining autonomous systems. The iBGP connections between routers D, E, and F provide consistent routing information to the autonomous systems.

Stub and multihomed autonomous systems

As shown in the preceding figure, [BGP networks](#) on page 14, an AS can include one or more BGP speakers that establish peer-to-peer sessions with BGP speakers in other autonomous systems to provide external route information for the networks within the AS.

A stub AS has a single BGP speaker that establishes a peer-to-peer session with one external BGP speaker. In this case, the BGP speaker provides external route information only for the networks within its own AS.

A multihomed AS has multiple BGP speakers.

Peers

BGP uses Transmission Control Protocol (TCP) as a transport protocol. When two routers open a TCP connection to each other for the purpose of exchanging routing information, they form a peer-to-peer relationship. In the preceding figure, [BGP networks](#) on page 14, Routers A and D are BGP peers, as are Routers B and E, C and E, F and G, and Routers D, E, and F.

Although Routers A and D run eBGP, Routers D, E, and F within AS 40 run iBGP. The eBGP peers directly connect to each other, while the iBGP peers do not. As long as an IGP operates and allows two neighbors to logically communicate, the iBGP peers do not require a direct connection.

* Note:

You cannot create the same iBGP peers on two different VRFs, or the same eBGP peers on two different chassis. Only one local autonomous system (AS) can exist for each chassis or VRF.

Because all BGP speakers within an AS must be fully meshed logically, the iBGP mesh can grow to large proportions and become difficult to manage. You can reduce the number of peers within an AS by creating confederations and route reflectors.

BGP peers exchange complete routing information only after the peers establish a connection. Thereafter, BGP peers exchange routing updates. An update message consists of a network number, a list of autonomous systems that the routing information passed through (the AS path), and other path attributes that describe the route to a set of destination networks. When multiple paths exist, BGP compares the path attributes to choose the preferred path. Even if you disable BGP, the system logs all BGP peer connection requests. For more information about update messages, see [BGP updates](#) on page 31.

Supernet advertisements

BGP has no concept of address classes. Each network listed in the network layer reachability information (NLRI) portion of an update message contains a prefix length field, which describes the length of the mask associated with the network. The prefix length field allows for both supernet and subnet advertisement. The supernet advertisement is what makes classless interdomain routing (CIDR) possible (see [CIDR and aggregate addresses](#) on page 19).

Bandwidth and maintenance reduction

BGP provides two features that reduce the high bandwidth and maintenance costs associated with a large full-mesh topology:

- confederations
- route reflectors

For information on confederations and route reflectors, see [Routing information consolidation](#) on page 19.

BGP 4 Byte AS Support

Each Autonomous System (AS) must have its own unique number. Because the 2-byte AS numbering scheme is unable to meet the increasing demand, the switch supports 4-byte AS numbers. This feature is enabled by supporting RFC 4893, BGP Support for 4-octet AS Number Space.

The switch supports the following three types of peer relationships as a result of 4 byte AS support:

- Old peer to old peer
- Old peer to new peer
- New peer to new peer

An old peer is the one that supports 2-byte AS numbers only and new peer is the one that supports both 2-byte AS numbers and 4-byte AS numbers.

RFC4893 supports two new path attributes:

- AS4_PATH contains the AS path encoded with a 4-octet AS number.
- AS4-AGGR is a new aggregator attribute that carries a 4-octet AS number.

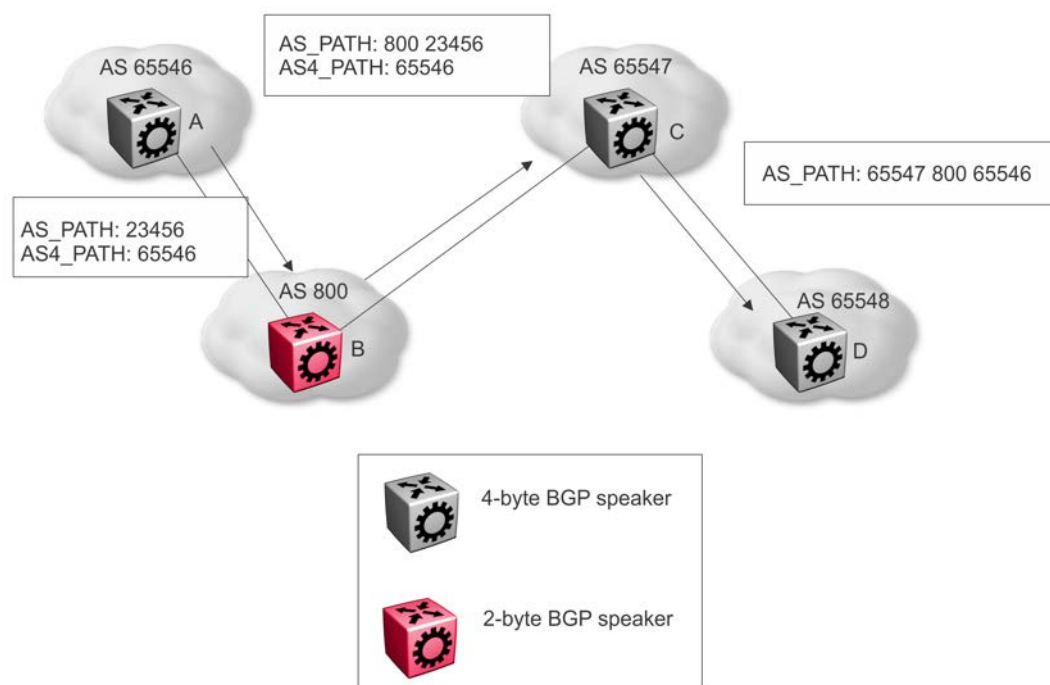


Figure 3: 2-byte and 4-byte Mixed Environment

The preceding figure shows an example of how the switch uses the AS4_PATH attribute in a mixed environment. The figure illustrates how a 2-byte BGP speaker interoperates with a 4-byte BGP speaker.

Router B is a 2-byte BGP speaker. Router A substitutes AS_PATH with the AS_TRANS, a 2-octet AS number defined by RFC4893 for backward compatibility, and encodes the 4-byte AS into AS4_PATH in BGP updates it sends to router B.

Router B does not understand the AS4_PATH but does preserve the information and sends it to router C.

Router C is a 4-byte BGP speaker. Router C merges the information received in AS_PATH and AS4_PATH, and encodes the 4-byte AS when it sends the AS_PATH information to router D.

Old Peer to Old Peer

When the peer relationship between an old peer and another old peer is established, 4 byte AS numbers contained in the AS4_PATH and AS4_AGGREGATOR are transited to other peers.

! Important:

Do not assign 23456 as an AS number. The Internet Assigned Numbers Authority (IANA) reserved this number for the AS_TRANS attribute and BGP uses it to facilitate communication between peer modes. AS_TRANS uses a 2-byte AS format to represent a 4-byte AS number. The switch interprets the AS_TRANS attribute and propagates it to other peers.

New Peer to New Peer

The new BGP speaker establishes its 4 byte AS support through BGP capability advertisement. A BGP speaker that announces such capability and receives it from its peer, uses 4 byte AS numbers

in AS_PATH and AGGREGATOR attributes and assumes these attributes received from its peer are encoded in 4 byte AS numbers.

The new BGP attributes AS4_PATH and AS4_AGGREGATOR received from the new BGP speaker between the new BGP peers in the update message is discarded.

Old Peer to New Peer

An old BGP speaker and a new BGP speaker can form peering relationship only if the new BGP speaker is assigned a 2 byte AS number. This 2 byte number can be any global unique AS number or AS_TRANS.

New BGP speaker sends AS path information to the old BGP speaker in AS_PATH attribute as well as AS4_PATH attribute. If the entire AS_PATH consists of only 2 byte AS numbers then the new BGP speaker does not send AS4_PATH information.

The 4-byte AS number feature does not in any way restrict the use or change the way you configure 2-byte AS numbers. You can also configure 2-byte AS or 4 byte AS numbers in AS path lists, community lists, and route policies.

BGP 4-byte AS Number Notation

BGP 4-byte AS numbers are represented in two ways: AS Plain and AS dot. The default form of representing the AS numbers is AS Plain while you have an option to configure AS dot. AS Plain form of representation is preferred over AS dot representation as a large amount of network providers find the AS dot notation incompatible with the regular expressions used by them. In case of any issues, troubleshooting and analyzing also gets difficult with AS dot notation.

BGP AS Number Format – AS Plain

Table 3: Default Asplain 4-Byte Autonomous System Number Format

Format	Configuration Format	Show Command Output and Regular Expression Match Format
asplain	2-byte: 1 to 65535	2-byte: 1 to 65535
	4-byte: 65536 to 4294967295	4-byte: 65536 to 4294967295
asdot	2-byte: 1 to 65535	2-byte: 1 to 65535
	4-byte: 1.0 to 65535.65535	4-byte: 65536 to 4294967295

BGP AS Number Format - ASdot

Table 4: Asdot 4-Byte Autonomous System Number Format

Format	Configuration Format	Show command output and Regular Expression Match Format
asplain	2-byte: 1 to 65535	2-byte: 1 to 65535
	4-byte: 65536 to 4294967295	4-byte: 1.0 to 65535.65535
asdot	2-byte: 1 to 65535	2-byte: 1 to 65535
	4-byte: 1.0 to 65535.65535	4-byte: 1.0 to 65535.65535

For more information on configuring 4 byte AS numbers, see [Configuring 4 byte AS numbers](#) on page 53.

Routing information consolidation

Use the information in this section to understand how to reduce the size of routing tables.

CIDR and aggregate addresses

Classless interdomain routing (CIDR) is an addressing scheme (also known as supernetting) that eliminates the concept of classifying networks into class types. Earlier addressing schemes identified five classes of networks: Class A, Class B, Class C, Class D, and Class E. This document does not discuss Classes D (used for multicast) and E (reserved and currently not used).

Network 195.215.0.0, an illegal Class C network number, becomes a legal supernet when represented in CIDR notation as 195.215.0.0/16. The /16 is the prefix length and expresses the explicit mask that CIDR requires. In this case, the addition of the prefix /16 indicates that the subnet mask consists of 16 bits (counting from the left).

Using this method, supernet 195.215.0.0/16 represents 195.215.0.0 255.255.0.0. The following table shows the conversion of prefix length to subnet mask.

Table 5: CIDR conversion

Prefix	Dotted-decimal	Binary	Network class
/1	128.0.0.0	1000 0000 0000 0000 0000 0000 0000 0000	128 Class A
/2	192.0.0.0	1100 0000 0000 0000 0000 0000 0000 0000	64 Class A
/3	224.0.0.0	1110 0000 0000 0000 0000 0000 0000 0000	32 Class A
/4	240.0.0.0	1111 0000 0000 0000 0000 0000 0000 0000	16 Class A
/5	248.0.0.0	1111 1000 0000 0000 0000 0000 0000 0000	8 Class A
/6	252.0.0.0	1111 1100 0000 0000 0000 0000 0000 0000	4 Class A
/7	254.0.0.0	1111 1110 0000 0000 0000 0000 0000 0000	2 Class A
/8	255.0.0.0	1111 1111 0000 0000 0000 0000 0000 0000	1 Class A or 256 Class B
/9	255.128.0.0	1111 1111 1000 0000 0000 0000 0000 0000	128 Class B
/10	255.192.0.0	1111 1111 1100 0000 0000 0000 0000 0000	64 Class B
/11	255.224.0.0	1111 1111 1110 0000 0000 0000 0000 0000	32 Class B
/12	255.240.0.0	1111 1111 1111 0000 0000 0000 0000 0000	16 Class B
/13	255.248.0.0	1111 1111 1111 1000 0000 0000 0000 0000	8 Class B
/14	255.252.0.0	1111 1111 1111 1100 0000 0000 0000 0000	4 Class B
/15	255.254.0.0	1111 1111 1111 1110 0000 0000 0000 0000	2 Class B

Table continues...

Prefix	Dotted-decimal	Binary	Network class
/16	255.255.0.0	1111 1111 1111 1111 0000 0000 0000 0000	1 Class B or 256 Class C
/17	255.255.128.0	1111 1111 1111 1111 1000 0000 0000 0000	128 Class C
/18	255.255.192.0	1111 1111 1111 1111 1100 0000 0000 0000	64 Class C
/19	255.255.224.0	1111 1111 1111 1111 1110 0000 0000 0000	32 Class C
/20	255.255.240.0	1111 1111 1111 1111 1111 0000 0000 0000	16 Class C
/21	255.255.248.0	1111 1111 1111 1111 1111 1000 0000 0000	8 Class C
/22	255.255.252.0	1111 1111 1111 1111 1111 1100 0000 0000	4 Class C
/23	255.255.254.0	1111 1111 1111 1111 1111 1110 0000 0000	2 Class C
/24	255.255.255.0	1111 1111 1111 1111 1111 1111 0000 0000	1 Class C

Use CIDR to assign network prefixes of arbitrary lengths, as opposed to the obsolete class system, which assigned prefixes as even multiples of an octet.

For example, you can assign a single routing table supernet entry of 195.215.16/21 to represent 8 separate Class C network numbers: 195.215.16.0 through 195.215.23.0.

Supernet addressing

You can create a supernet address that covers an address range.

For example, to create a supernet address that covers an address range of 192.32.0.0 to 192.32.9.255, perform the following steps:

1. Convert the starting and ending address range from dotted-decimal notation to binary notation (see the following figure).

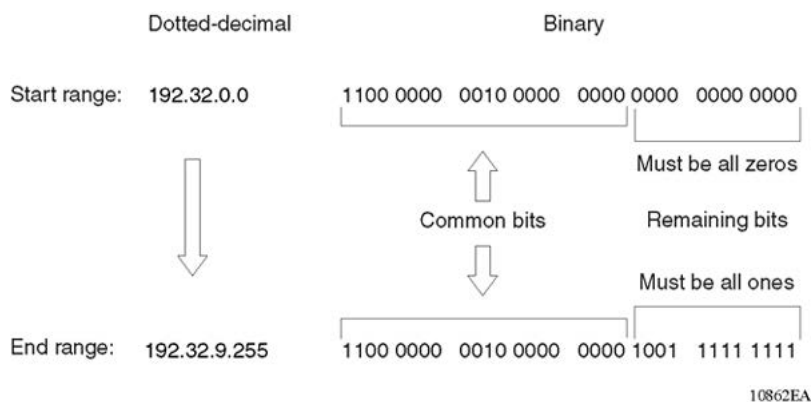
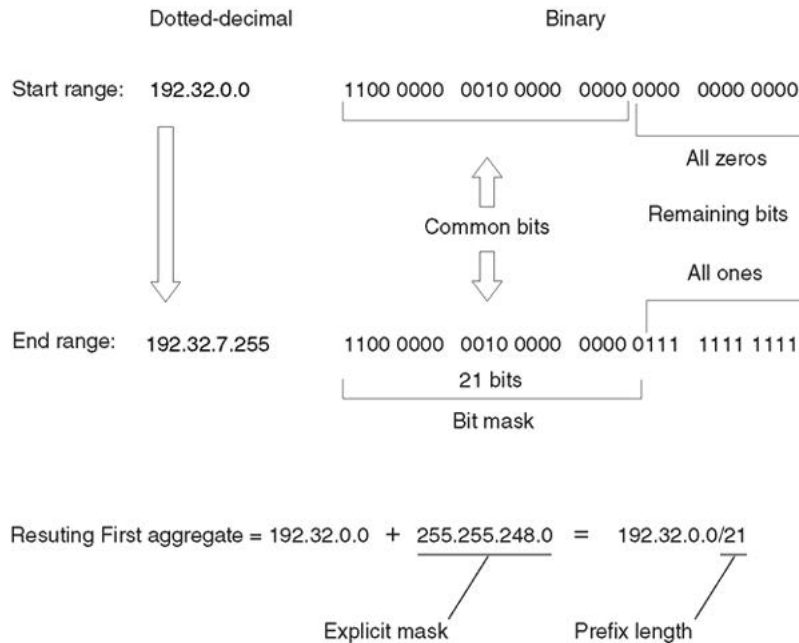


Figure 4: Binary notation conversion

2. Locate the common bits in both ranges. Ensure that the remaining bits in the start range are zeros, and the remaining bits in the end range are all ones. In this example, the remaining bits in the end range are not all ones.
3. If the remaining bits in the end range are not all ones, you must recalculate to find the IP prefix that has only ones in the remaining bits in the end range.
4. Recalculate to find a network prefix that has all ones in the remaining end range bits (see the following figure). In this example, 192.32.7.255 is the closest IP prefix that matches the common bits for the start range.



10863EA

Figure 5: First aggregate and prefix length

5. The 21 bits that match the common bits form the prefix length. The prefix length is the number of binary bits that form the explicit mask (in dotted-decimal notation) for this IP prefix.
6. The remaining aggregate is formed from 192.32.8.0 to the end range, 192.32.9.255.

As shown in [Figure 5: First aggregate and prefix length](#) on page 21, the resulting first aggregate 192.32.0.0/21 represents all of the IP prefixes from 192.32.0.0 to 192.32.7.255.

The following figure shows the results after forming the remaining aggregate from 192.32.9.0 to the end range, 192.32.9.255.

The resulting aggregate 192.32.8.0/23 represents all of the IP prefixes from 192.32.8.0 to 192.32.9.255.

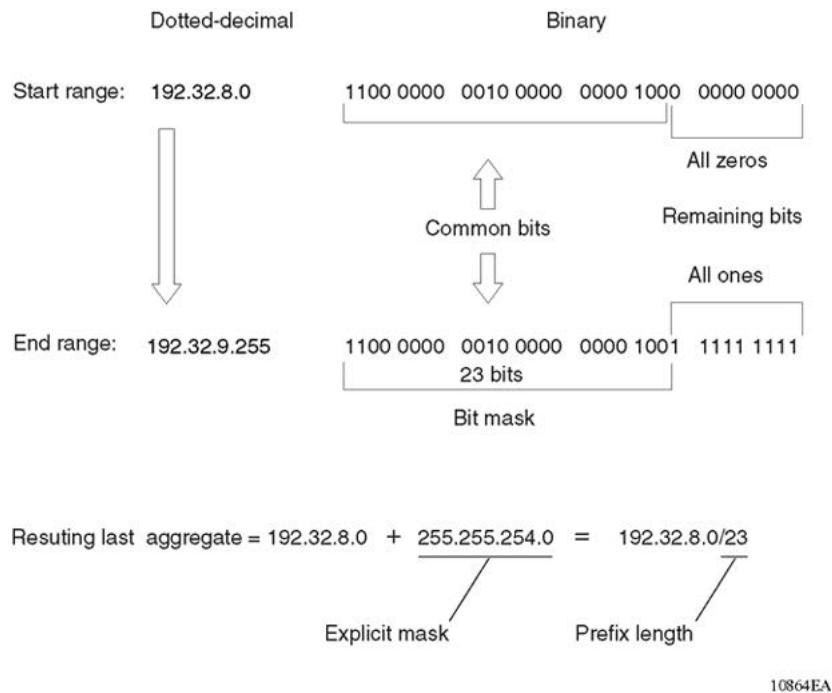


Figure 6: Last aggregate and prefix length

The final result of calculating the supernet address that ranges from 192.32.0.0 to 192.32.9.255 is as follows:

192.32.0.0 (with mask) 255.255.248.0 = 192.32.0.0/21

192.32.8.0 (with mask) 255.255.254.0 = 192.32.8.0/23

Aggregate routes

Eliminating the idea of network classes provides an easy method to aggregate routes. Rather than advertise a separate route for each destination network in a supernet, BGP uses a supernet address to advertise a single route (called an aggregate route) that represents all the destinations. CIDR also reduces the size of the routing tables used to store advertised IP routes.

The following figure shows an example of route aggregation using CIDR. In this example, a single supernet address 195.215.0.0/16 advertises 256 separate Class C network numbers 195.215.0.0 through 195.215.255.0.

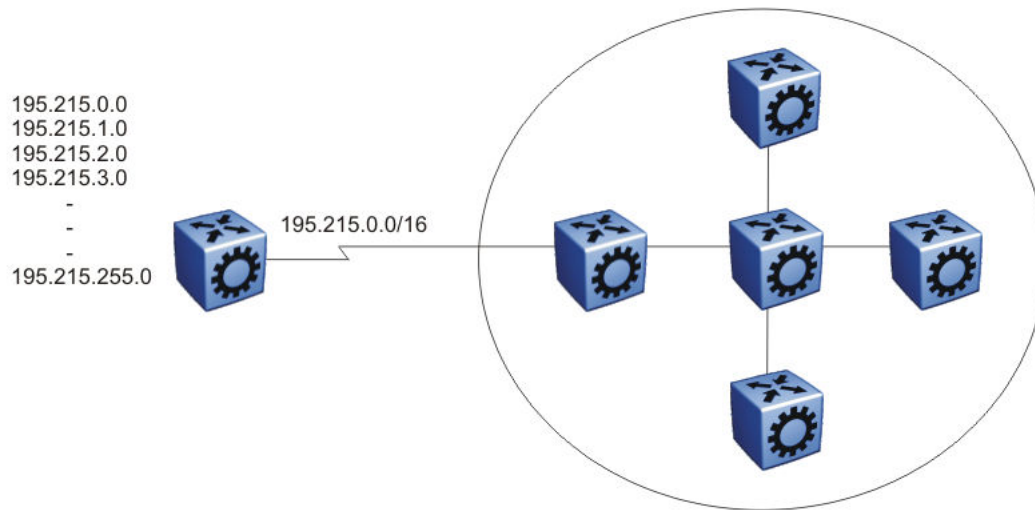


Figure 7: Aggregating routes with CIDR

Confederations

A BGP router configured for iBGP establishes a peer-to-peer session with every other iBGP speaker in the AS. In an AS with a large number of iBGP speakers, this full-mesh topology can result in high bandwidth and maintenance costs.

As shown in the following example, a full-mesh topology for an AS with 50 iBGP speakers requires 1225 internal peer-to-peer connections:

Example:

$$n \times (n-1)/2 = n \text{ iBGP sessions}$$

where:

$$50 \times (50-1)/2 = 1225 \text{ number of unique iBGP sessions}$$

You can reduce the high bandwidth and maintenance costs associated with a large full-mesh topology by dividing the AS into multiple smaller autonomous systems (sub-autonomous systems), and then group them into a single confederation (see the following figure).

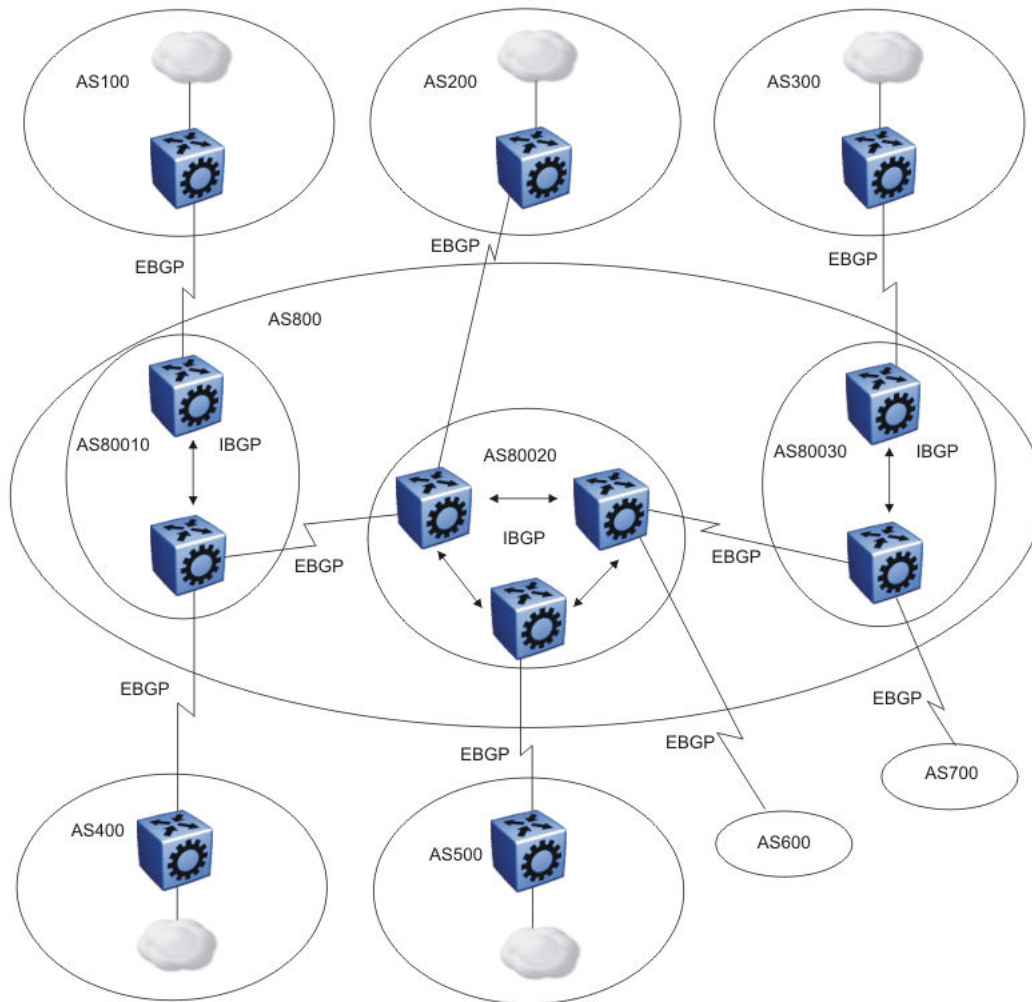


Figure 8: Confederations

As shown in the preceding figure, each sub-AS is fully meshed within itself and has eBGP sessions with other sub-autonomous systems in the same confederation.

Although the peers in different autonomous systems have eBGP sessions with the various sub-AS peers, they preserve the next-hop, Multi-Exit Discriminator (MED), and local preference information and exchange routing updates as if they were iBGP peers. All of the autonomous systems retain a single interior gateway protocol (IGP). When the confederation uses its own confederation identifier, the group of sub-autonomous systems appear as a single AS (with the confederation identifier as the AS number).

Route reflectors

Another way to reduce the iBGP mesh inherent in an AS with a large number of iBGP speakers is to configure a route reflector. Using this method, when an iBGP speaker needs to communicate with other BGP speakers in the AS, the speaker establishes a single peer-to-peer route reflector client session with the iBGP route reflector.

In an AS, more than one route reflector cluster can exist and more than one route reflector in a cluster. When more than one reflector exists in a cluster, take care to prevent route loops.

The following figure shows a simple iBGP configuration with three iBGP speakers (routers A, B, and C). Without route reflectors, after Router A receives an advertised route from an external neighbor, it must advertise the route to Routers B and C.

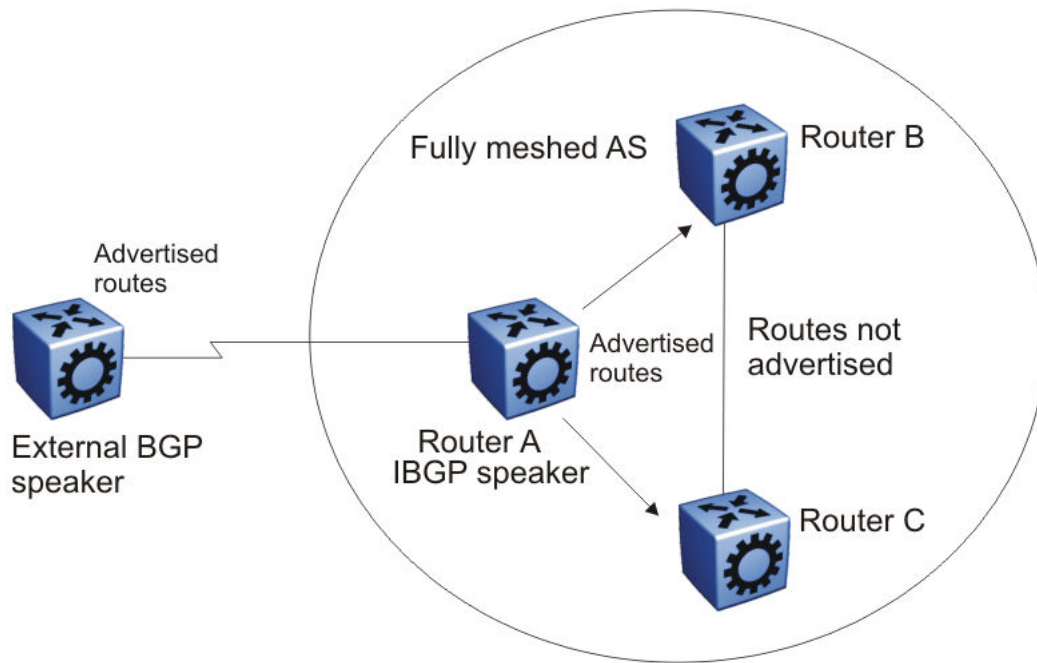


Figure 9: Fully meshed AS with iBGP speakers

Routers B and C do not readvertise the iBGP learned routes to other iBGP speakers. BGP does not allow routers to pass routes learned from internal neighbors on to other internal neighbors, which avoids routing information loops.

As shown in the following figure, when you configure an internal BGP peer (Router B) as a route reflector, all of the iBGP speakers do not need to be fully meshed. In this case, the assigned route reflector passes iBGP learned routes to a set of iBGP neighbors.

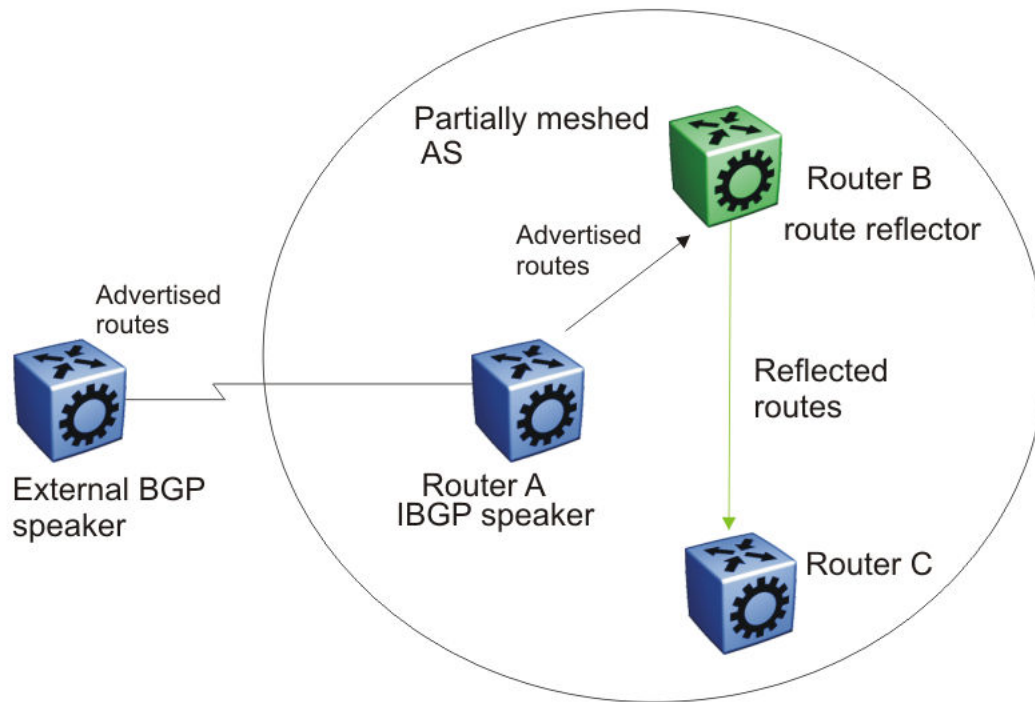


Figure 10: AS with route reflector

After Router B, the route reflector, receives routes from Router A (the iBGP speaker), it advertises them to router C. Conversely, after the route reflector receives routes from internal peers, it advertises those routes to Router A. Routers A and C do not need an iBGP session.

Route reflectors separate internal peers into two groups: client peers and nonclient peers. The route reflector and its clients form a cluster. The client peers in the cluster do not need to be fully meshed, and do not communicate with iBGP speakers outside their cluster. Nonclient peers must be fully meshed with each other.

The following figure shows a cluster, where Router A is the route reflector in a cluster with client routers B, C, and D. Routers E, F, and G are fully meshed, nonclient routers.

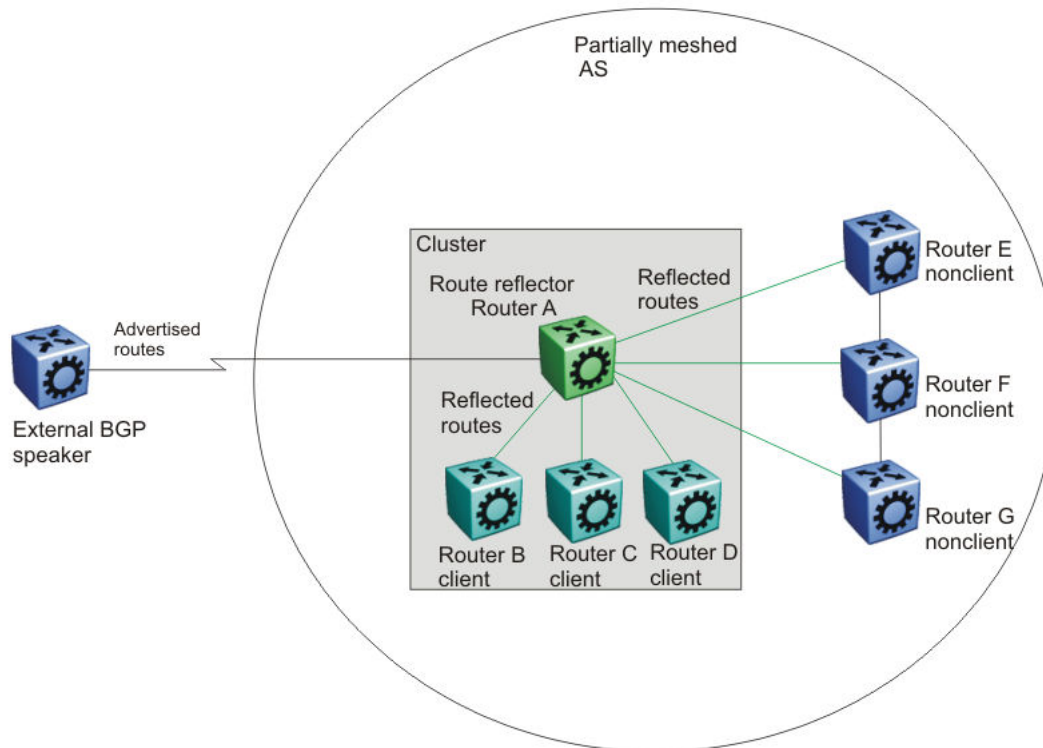


Figure 11: Route reflector with client and nonclient peers

BGP communities

You can group destinations into communities to simplify policy administration. A community is a group of destinations that share a common administrative property.

Use community control routing policies with respect to destinations. Create communities when you have more than one destination and want to share a common attribute.

The following list identifies specific community types:

- Internet—advertise this route to the Internet community
- no advertise—do not advertise to BGP peers including iBGP peers

You can use a community to control which routing information to accept, prefer, or distribute to other BGP neighbors. If you specify the append option in the route policy, the router adds the specified community value to the existing value of the community attribute. Otherwise, the specified community value replaces a previous community value.

BGP path attributes

You can create policies that control routes, work with default routing, control specific and aggregated routes, and manipulate BGP path attributes.

Four categories of BGP path attributes exist:

- Well-known mandatory attributes must be in every BGP update message.
- Well-known discretionary attributes can be in a BGP update message.
- Optional transitive attributes are accepted and passed to other BGP peers.
- Optional non-transitive attributes can be either accepted or ignored, but must not pass along to other BGP peers.

Border routers that utilize built-in algorithms or manually configured policies to select paths use path attributes. BGP uses the following path attributes to control the path a BGP router chooses:

- origin (well-known mandatory)
- AS_path (well-known mandatory)
- next hop (well-known mandatory)
- MED attribute (optional non-transitive)
- local preference (well-known discretionary)
- atomic aggregate (well-known discretionary)
- aggregator (optional transitive)
- community (optional transitive)

For more information about path attributes in BGP updates, see [Path Attributes](#) on page 32.

BGP Route Selection

A BGP router determines the best path to a destination network. This path is then eligible for use in the IP forwarding table and the router also advertises the path to its eBGP peers. To choose the best of multiple BGP routes to a destination, the router executes a best path algorithm.

The algorithm chooses a route in the following order:

- highest weight

Weight is a locally significant parameter associated with each BGP peer. You can use the weight to influence which peer paths the router uses.

- highest local preference

The local preference has global significance within an AS. You can manipulate the preference using route policies to influence path selection.

- prefer locally originated paths

The router prefers a path locally originated using the network, redistribution, or aggregate command over a path learned through a BGP update. The router prefers local paths sourced by network or redistribute commands over local aggregates sourced by the aggregate address command.

- shortest AS path

The AS path parameter specifies the autonomous systems that the network prefix traversed. The AS path commonly determines the best path. For example, a router can choose a path based on whether the network passed through a specific AS. You can configure a route policy to match the AS, and then modify the local preference. Also, you can pad the AS path before the AS advertises it to a peer AS, so that downstream routers are less likely to prefer the advertised network path.

The AS_CONFED_SEQUENCE length will also be considered while picking the best path inside the confederation.

- lowest origin type

The order of preference is IGP, EGP, INC (incomplete).

- lowest MED

The MED parameter influences the preferred path from a remote AS to the advertising AS. This parameter applies when there are multiple exit points from the remote AS to the advertising AS. A lower MED value indicates a stronger path preference than a higher MED value. By default, the MED attribute is ignored as specified by the BGP global parameter Always Compare MED except when the routes come from the same AS. This parameter must be enabled for MEDs to be compared (and for this step of the best path algorithm to execute).

The router compares MEDs regardless of what the first (neighboring) AS specified in the AS_PATH. Deterministic MED, when enabled, means that the first AS of the multiple paths must be the same. Paths received with no MED are assigned a MED of 0, unless the global BGP parameter Missing Is Worst is enabled. If so, received paths are assigned a MED of 4 294 967 294. Missing Is Worst is enabled by default. The "no-med-path-is-worst" flag has an impact only when the "First AS" or the "Most Left AS" is the same for multiple routes received. The router changes paths received with a MED of 4 294 967 295 to 4 294 967 294 before insertion into the BGP table.

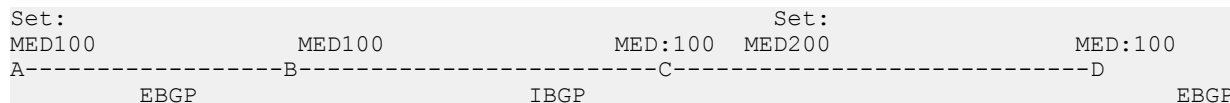
 Note:

You cannot enable or disable the MED selection process. BGP aggregation does not occur when routes have different MEDs or next hops.

When MED value is set in route-map configuration, the configured MED value is not applicable if it is already set in the associated Path Attribute.

1. When router A sets MED value of 100 by route-map, it will send Path Attribute with MED=100 to EBGp peer B.
2. Router B sends Path Attribute with MED=100 to IBGP peer C.

3. If the route-map is configured with "set MED 200", then router C does not apply MED=200 to the Path Attribute as it is already set to 100 when it is received from router B.
4. Router D will get Path Attribute with MED=100 so that router C does not influence router D when it selects the best route.



Example: If Prefix: X is set as MED=100 from router A, it will be received at B with MED=100, and will carry same MED=100 value to router C, as it is an iBGP peer. Router C will not propagate MED=100 value to D as MED is a non-transitive attribute, so MED can travel maximum of 1 AS.

- lowest IGP metric to the BGP next-hop

If multiple paths exist whose BGP next-hop is reachable through an IGP, the path with the lowest IGP metric to the BGP next-hop is chosen.

- prefer external paths (learned by eBGP) over internal paths (iBGP)

The system prefers external paths over internal paths.

- if Equal Cost Multipath (ECMP) is enabled, insert up to four paths in the routing table

If you enable ECMP, multiple BGP learned routes that use the same metric to different IP next-hops are installed in the IP forwarding table for traffic load-balancing purposes.

- lowest router ID

The lowest router ID, or Circuitless IP (CLIP) address, is preferred.

BGP and dampened routes

The switch supports route dampening (route suppression). When you use route dampening, a route accumulates penalties each time the route fails. After the accumulated penalties exceed a threshold, the router no longer advertises the route. The router enters the suppressed routes into the routing table only after the accumulated penalty falls below the reuse threshold.

Route flap dampening suppresses the advertisement of the unstable route until the route becomes stable. For information about how to enable flap-dampening, see [Configuring BGP globally](#) on page 48. For information about viewing flap dampening configurations, see [Viewing global flap-dampening configurations](#) on page 81.

Dampening applies only to routes that are learned through an eBGP. Route flap dampening prevents routing loops and protects iBGP peers from having higher penalties for routes external to the AS.

The following paragraph describes the algorithm that controls route flaps.

After the route flaps the first time

- the router creates a route history entry
- a timer starts (180 seconds)

If the route does not flap again, the router uses this timer to delete the history entry after the 180 seconds expires.

After the route flaps a second time

- The penalty is recalculated based on the decay function.

If the penalty is greater than the cut-off value (1536), the route is suppressed and the reuse time is calculated based on the reuse time function.

- The reuse timer starts.

After the reuse time expires, the suppressed route is announced again (the reuse time is recalculated if the route flaps again). The penalty decays slower for withdrawn routes than for update routes. The route history entry is kept longer if the route is withdrawn. For update history, the delete time is 90 seconds and the withdrawn history delete time is 180 seconds.

BGP Updates

BGP uses update messages to communicate information between two BGP speakers. The update message can advertise a single feasible route to a peer, or withdraw multiple unfeasible routes from service.

The following figure shows the format of an update message.

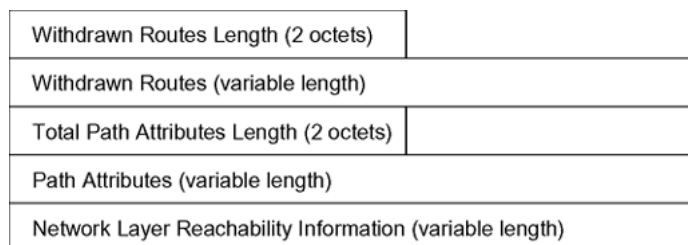


Figure 12: Update Message Format

This section describes how BGP uses the update message fields to communicate information between BGP speakers.

Withdrawn Routes Length

The withdrawn routes length parameter (referred to in RFC1771 as the Unfeasible Routes Length field) indicates the total length of the withdrawn routes field in octets. The withdrawn routes length field calculates the length of the NLRI field. For example, a value of 0 indicates that no routes are withdrawn from service, and that the withdrawn routes field is not present in this update message.

Withdrawn Routes

The withdrawn routes parameter is a variable-length parameter that contains a list of IP prefixes for routes that are withdrawn from service. The following figure shows the format of an IP prefix.



Figure 13: IP Prefix Format

The length indicates the number of bits in the prefix (also called the network mask).

For example, 192.0.2.0/24 is equivalent to 192.0.2.0 255.255.255.0 (the /24 indicates the number of bits in the length parameter to represent the network mask 255.255.255.0).

The prefix parameter contains the IP address prefix itself, followed by enough trailing bits to make the length of the whole field an integer multiple of 8 bits (1 octet).

Total Path Attributes Length

The total path attributes length parameter indicates the total length of the path attributes parameter in octets.

The total path attributes length calculates the length of the NLRI parameter. For example, a value of 0 indicates that no NLRI field is present in this update message.

Path Attributes

The path attributes parameter is a variable-length sequence of path attributes that exists in every BGP update. The path attributes contain BGP attributes associated with the prefixes in the NLRI parameter.

For example, the attribute values allow you to specify the prefixes that the BGP session can exchange, or which of the multiple paths of a specified prefix to use.

The attributes carry the following information about the associated prefixes:

- the path origin
- the AS paths through which the prefix is advertised
- the metrics that display degrees of preference for this prefix

The following figure shows the encoding used with the path attribute parameter.

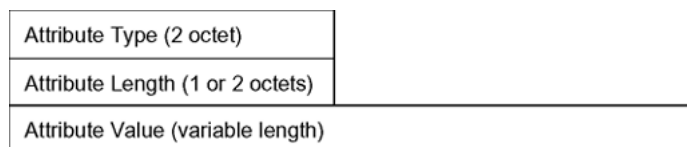


Figure 14: Path Attribute Encoding

Attribute Type

As shown in the following figure, the attribute type is a two-octet field that comprises two sub-fields: attribute flags and attribute type code.

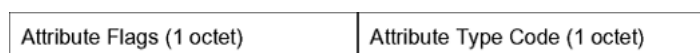


Figure 15: Attribute Type Fields

The attribute flags parameter is a bit string that contains four binary values that describe the attribute, and four unused bits. The following list provides bit descriptions (from the high-order bit to the low-order bit):

- The high-order bit (bit 0) is the optional bit. When this bit is set (the value is 1), the attribute is optional. When this bit is clear (the value is 0), the attribute is well-known. Well-known attributes must be recognized by all BGP implementations and, when appropriate, passed on to BGP peers. Optional attributes are not required in all BGP implementations.
- The second high-order bit (bit 1) is the transitive bit. For well-known attributes, this bit must be set to 1. For optional attributes, it defines whether the attribute is transitive (when set to 1) or non-transitive (when set to 0).
- The third high-order bit (bit 2) is the partial bit. The partial bit defines whether the information in the optional transitive attribute is partial (when set to 1) or complete (when set to 0). For well-known attributes and for optional non-transitive attributes the partial bit must be set to 0.
- The fourth high-order bit (bit 3) is the extended length bit. The extended length bit defines whether the attribute length is one octet (when set to 0) or two octets (when set to 1). The attribute flag can use the extended length only if the length of the attribute value is greater than 255 octets.
 - If the extended length bit of the attribute flags octet is set to 0, the third octet of the path attribute contains the length of the attribute data in octets.
 - If the extended length bit of the attribute flags octet is set to 1, then the third and the fourth octets of the path attribute contain the length of the attribute data in octets.
- The lower-order four bits of the attribute flags octet are unused. The lower-order four bits must be zero (and must be ignored when received).

The attribute type code parameter contains the attribute type code, as defined by the Internet Assigned Numbers Authority (IANA). The attribute type code uniquely identifies the attribute from all others. The remaining octets of the path attribute represent the attribute value and are interpreted according to the attribute flags and the attribute type code parameters.

The following table shows the supported attribute type codes.

Table 6: BGP Mandatory Path Attributes

Attribute	Type code	Description
Origin	1	Defines the origin of the path information: • Value = 0 --- IGP (the path is valid all the way to the IGP of the originating AS) • Value = 1--- EGP (the last AS in the AS path uses an EGP to advertise the path) • Value = 2--- Incomplete (the path is valid only to the last AS in the AS path)

Table continues...

Attribute	Type code	Description
AS path	2	Contains a list of the autonomous systems that packets must traverse to reach the destinations. This code represents each AS path segment as follows: • path segment type • path segment length • path segment value
Next hop	3	Specifies the IP address of the border router to use as a next hop for the advertised destinations (destinations listed in the NLRI field of the update message).
Multixit discriminator	4	Discriminates among multiple exit or entry points to the same neighboring AS on external (internal-AS) links.
Local preference	5	Indicates the preference that AS border routers assign to a chosen route when they advertise it to iBGP peers
Atomic aggregate	6	Ensures that certain NLRI is not deaggregated
Aggregator	7	Identifies which AS performed the most recent route aggregation. This attribute contains the last AS number that formed the aggregate route followed by the IP address of the BGP speaker that formed the aggregate route.

Attribute Length

The attribute length can be one or two octets in length, depending on the value of the extended length parameter in the attributes flag field.

This parameter indicates the length of the attribute value field.

Attribute Value

The attribute value contains the actual value of the specific attribute. The system implements the attribute value according to the values in the attribute flags and the attribute type code parameters.

NLRI

The NLRI parameter is a variable length field that contains a list of prefixes. The packet size that BGP speakers can exchange limits the number of prefixes in the list.

Equal Cost Multipath

Equal Cost Multipath (ECMP) support allows a BGP speaker to perform route or traffic balancing within an AS by using multiple equal-cost routes submitted to the routing table by OSPF, RIP, or static routes.

For more information about ECMP, see [Configuring IPv4 Routing for VOSS](#).

MD5 message authentication

Authenticate BGP messages by using Message Digest 5 (MD5) signatures. After you enable BGP authentication, the BGP speaker verifies that the BGP messages it receives from its peers are actually from a peer and not from a third party masquerading as a peer.

BGPv4 TCP MD5 message authentication provides the following features:

- A TCP MD5 signature can exist for BGP peers. You can configure authentication and secret keys for each peer. Peers configured with common secret keys can authenticate each other and exchange routing information.
- The switch can concurrently have BGP peers with authentication enabled and other BGP peers with authentication disabled.
- The switch always encrypts the secret keys.

After you enable BGPv4 TCP MD5 authentication, the router computes an MD5 signature for each TCP packet based on the TCP packet and an individual peer secret key. The router adds this MD5 signature to the TCP packet that contains a BGP message and sends it with the packet, but it does not send the secret key.

The receiver of the TCP packet also knows the secret key and can verify the MD5 signature. A third party that tries to masquerade as the sender, however, cannot generate an authentic signature because it does not know the secret key.

In commands, the term password refers to the secret key. The secret keys provide security. If the keys are compromised, then the authentication itself is compromised. To prevent this, the switch stores the secret keys in encrypted form.

MD5 signature generation

BGP peers calculate MD5 signatures in BGP messages based on the following elements:

- TCP pseudo-header
- TCP header, excluding options
- TCP segment data
- TCP MD5 authentication key

If TCP receives an MD5 authentication key, it reduces its maximum segment size by 18 octets, which is the length of the TCP MD5 option. TCP adds an MD5 signature to each transmitted packet. The peer inserts the resulting 16-byte MD5 signature into the following TCP options: kind=19, length=18.

MD5 signature verification

After the switch receives a packet, it performs three tests. The following table lists the tests and the event message that TCP logs if a test fails.

Table 7: MD5 signature verification rules on BGP TCP packets

Condition tested	Action on success	Failure event message
Is the connection configured for MD5 authentication?	Verify that the packet contains a kind=19 option.	TCP MD5 No Signature
Is MD5 authentication enabled for this TCP connection?	TCP computes the expected MD5 signature.	TCP MD5 Authentication Disabled
Does the computed MD5 signature match the received MD5 signature?	TCP sends the packet to BGP.	TCP MD5 Invalid Signature

If a packet passes a test, it proceeds to the next test. After a packet passes all three tests, TCP accepts the packet and sends it to BGP.

If a packet fails a test, the switch logs an event, increments the count of TCP connection errors (wfTcConnMd5Errors), and discards the packet. The TCP connection remains open.

BGP and route redistribution

Redistribution imports routes from one protocol to another. Redistribution sends route updates for a protocol-based route through another protocol. For example, if OSPF routes exist in a router and they must travel through a BGP network, then configure redistribution of OSPF routes through BGP. This sends OSPF routes to a router that uses BGP.

The switch can redistribute routes:

- on an interface basis.
- on a global basis between protocols on a single VRF instance (intraVRF).
- between the same or different protocols on different VRF instances (interVRF).

Configure interface-based redistribution by configuring a route policy and apply it to the interface. Configure the match parameter to the protocol from which to learn the routes.

You can redistribute routes on a global basis, rather than on an interface basis. Use the **ip bgp redistribute** command to accomplish the (intraVRF) redistribution of routes through BGP, so that BGP redistribution occurs globally on all BGP-enabled interfaces. This redistribution does not require a route policy, but you can use one for more control.

If you configure redistribution globally and on an interface, redistribution through the route policy takes precedence.

You can redistribute routes from a protocol in one VRF to BGP in another VRF. You can use a route policy for redistribution control. If you enable route redistribution between VRF instances, ensure that IP addresses do not overlap.

Use caution when you configure redistribution. An improperly configured parameter can cause the router to advertise learned eBGP routes out of your local AS. If this happens, the local AS can route other networks.

Do not use redistribution if you peer to an Internet Service Provider (ISP) and do not want traffic to transit your local AS.

When you redistribute OSPF routes into BGP, route priorities can create routing loops. Because BGP has a higher route preference than OSPF external type 1 and 2 routes, if you redistribute OSPF external type 1 and 2 routes into BGP, the router uses the BGP routes, which can cause a routing loop.

Route-maps and BGP neighbors

BGP Routing Information Base (BGP RIB) stores routing information received from different peers. BGP RIB has two types of BGP routes, External and Internal (Local). The routes learned from BGP neighbors are External routes and all imported routes are considered as Internal (Local) routes.

In BGP RIB, the OSPF routes redistributed into BGP are considered as Internal (Local) and are matched by route-type only when the keyword is set to `local`. When match route-type is set to `external`, the route-maps applied on BGP neighbors are ignored and the set operation is not performed.

* Note:

This is applied only on the route-maps applied to BGP neighbors in BGP RIB, and not considered when applying a route-map to the redistribute command.

BGP route redistribution and DvR

DvR Controllers redistribute routes (direct routes, static routes and the default route) into the DvR domain. You can configure redistribution of DvR host routes into BGP.

For information on DvR, see [Configuring IPv4 Routing for VOSS](#).

BGP+

The switch extends the BGPv4 process to support the exchange of IPv6 routes using BGPv4 peering. BGP+ is an extension of BGPv4 for IPv6, which is indicated using the Address Family Identifier (AFI) in the BGP header.

The switch supports capabilities for AFI with the following values: 1 (IPv4) and 2 (IPv6). If the switch receives an OPEN message advertising an AFI with a different value, the connection is closed and a BGP notification message is sent to the peer mentioning unsupported capability.

BGP+ is only supported on the global VRF instance.

* Note:

Ensure you configure IPv6 forwarding for BGP+ to work.

Note that the BGP+ support on the switch is not an implementation of BGPv6. Native BGPv6 peering uses the IPv6 Transport layer (TCPv6) for establishing the BGPv6 peering, route exchanges, and data traffic.

The switch supports the exchange of IPv6 reachability information over IPv4 transport. To support BGP+, the switch supports two BGP protocol extensions, standards RFC 4760 (multi-protocol extensions to BGP) and RFC 2545 (MP-BGP for IPv6). These extensions allow BGPv4 peering to be enabled with IPv6 address family capabilities.

The implementation of BGP+ on the switch uses an existing TCPv4 stack to establish a BGPv4 connection. Optionally, nontransitive BGP properties are used to transfer IPv6 routes over the BGPv4 connection. Any BGP+ speaker has to maintain at least one IPv4 address to establish a BGPv4 connection.

Different from IPv4, IPv6 introduces scoped unicast addresses, identifying whether the address is global or link-local. When BGP+ is used to convey IPv6 reachability information for interdomain routing, it is sometimes necessary to announce a next hop attribute that consists of a global address and a link-local address. For BGP+, no distinction is made between global and site-local addresses.

The BGP+ implementation includes support for BGPv6 policies, including redistributing BGPv6 into OSPFv3, ISIS, RIPng, and advertising OSPFv3, ISIS, RIPng, IPv6 static and local routes into BGPv6 (through BGP+). It also supports the aggregation of global unicast IPv6 addresses.

When configuring BGP+ on the router that is enabled only for IPv6 (the router does not have an IPv4 address), then BGP router ID must be manually configured for the router.

BGP+ does not support confederations. You can configure confederations for IPv4 routes only.

The basic configuration of BGP+ is the same as BGPv4 with one additional parameter added and some existing commands altered to support IPv6 capabilities. You can enable and disable IPv6 route exchange by specifying the address family attribute as IPv6. Note that an IPv6 tunnel is required for the flow of IPv6 data traffic.

BGP+ tunnel

When you use BGP+ you must configure an IPv6 tunnel and static routes at BGP+ peers.

When BGP+ peers advertise route information, they use Update messages to advertise route information.

These RTM routes contain next-hop addresses from the BGP peer that the route was learned from.

The static routes correlate the next-hop addresses represented by the IPv4-mapped IPv6 address to a specific outgoing interface.

Following is one way to express a static route in an IPv6-configured tunnel for BGP+:

```
ipv6 route 2001:DB8:0:0:0:ffff:192.0.2.0/24 cost 1 tunnel 10 where
2001:DB8:0:0:0:ffff:192.0.2.0 is the IPv4-mapped IPv6 address of the BGP peer at 192.0.2.0
```

ECMP with BGP+

The ECMP feature supports and complements BGP+ protocol.

The number of equal-cost-paths supported can differ by hardware platform. For more information, see [Release Notes for XA1400 Series](#).

You can use multiple paths for load sharing of traffic. These multiple paths allow faster convergence to other active paths in case of network failure. By maximizing load sharing among equal-cost paths, you can use your links between routers more efficiently when sending IP and IPv6 traffic. Equal Cost Multipath is formed using routes from the same protocol.

*** Note:**

To add BGP+ equal cost paths in the routing table, you must enable the following:

- IPv6 ECMP feature globally
- BGP multiple-paths attribute

BGPv6

BGP peering over IPv6 transport uses a BGPv6 peer to exchange IPv6 routes over an IPv6 transport layer. This is different than BGP+, which enables exchange of IPv6 routes over a BGPv4 peer. Also with BGP+, you must use an IPv6 tunnel to install and configure IPv6 routes in an IPv6 Routing Table Manager (RTM). BGP+ uses an IPv4 mapped IPv6 address for the next hop address and requires you to configure IPv6 static routes and install IPv6 routes in an IPv6 RTM where the next hop for the static route is an IPv6 tunnel interface.

BGPv6 supports the following:

- Input/Output policies.
- Redistribution of OSPFv3, IS-IS, IPv6 static route, and IPv6 direct routes into BGPv6.
- Aggregation of global unicast IPv6 addresses.

*** Note:**

BGP+ also supports the preceding features.

RFC

The switch supports the BGP multiprotocol extension, as described in RFC 4760. Also supports RFC 2545 (MP-BGP for IPv6).

The BGP protocol extensions ensure peering can be enabled with IPv6 address family capabilities.

Route exchange

BGPv6 does not exchange any IPv4 routes. BGPv6 advertises or learns only IPv6 routes.

The following table shows the differences between BGPv4 and BGPv6 for route exchange.

	GRT/VRF	IPv4 Routes Exchange	IPv6 Routes Exchange
BGPv4	GRT	Supported	Supported (BGP+)
	VRF	Supported	Not supported

Table continues...

	GRT/VRF	IPv4 Routes Exchange	IPv6 Routes Exchange
			 Note: IPv6 over IPv4 tunnels is not yet virtualized.
BGPv6	GRT	Not supported	Supported
	VRF	Not supported	Supported

Specify the address family attribute as IPv6 to enable IPv6 route exchange.

You can enable IPv6 route exchange by specifying the address family attribute as IPv6. Optionally, you can use non-transitive BGP properties to exchange IPv6 routes between the BGPv6 peering. Any BGPv6 speaker must maintain at least one IPv6 address to establish a BGPv6 connection. The IPv6 scoped unicast addresses can identify the address as global or link-local. If you use BGPv6 to convey IPv6 reachability information for interdomain routing, you can also announce a next hop attribute that consists of a global address and a link-local address.

 Note:

BGPv6 does not support adjacency on link-local.

Authentication

BGPv6 uses IPsec for security. MD-5 authentication is supported for BGPv4 and is not supported for BGPv6.

The following table shows the differences between BGPv4 and BGPv6 for authentication.

	MD5	IPsec	SHA1/SHA2
BGPv4	Supported	Not supported	Not supported
BGPv6	Not supported	Supported  Note: IP Sec is not virtualized, hence BGPv6 is supported only in Global Router mode, and not supported in VRF mode.	Not supported

MD5 authentication

MD5 authentication is not supported in BGPv6 so it is not necessary to enable MD5 authentication.

IPsec

Only IPsec is supported. Therefore, MD5 authentication cannot be configured.

Consistency checking

Includes consistency checking for MD5 authentication. BGP peer and BGP peer group configuration for IPv6 addresses include a rule to block MD5 authentication. If you attempt to configure MD5 authentication, you will receive an error message.

IPv6 tunneling

With BGPv6, IPv6 tunneling is not required for IPv6 data traffic flow. An IPv6 tunnel is required for BGP+.

Circuitless IP

Circuitless IP (CLIP) is a virtual (or loopback) interface that you do not associate with a physical port. You can use a CLIP interface to provide uninterrupted connectivity to your switch as long as an actual path exists to reach the device. For example, as shown in the following figure, a physical point-to-point link exists between R1 and R2 along with the associated addresses (195.39.1.1/30 and 195.39.1.2/30). Note also that an iBGP session exists between two additional addresses 195.39.128.1/32 (CLIP 1) and 195.39.128.2/32 (CLIP 2).

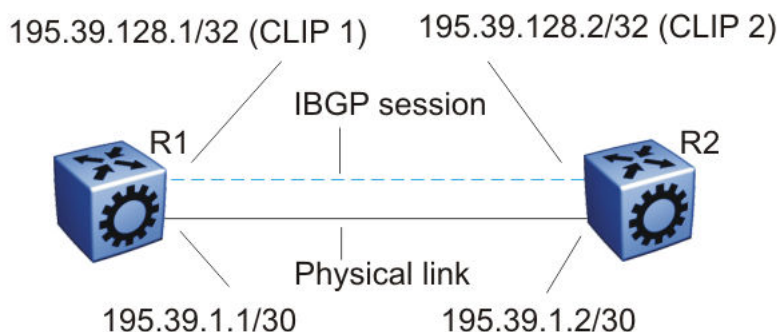


Figure 16: Routers with iBGP connections

The system treats the CLIP interface like an IP interface and treats the network associated with the CLIP as a local network attached to the device. This route always exists and the circuit is always up because no physical attachment exists.

The router advertises routes to other routers in the domain either as external routes using the route-redistribution process or after you enable OSPF in a passive mode to advertise an OSPF internal route. You can configure only the OSPF protocol on the CLIP interface. After you create a CLIP interface, the system software programs a local route with the CPU as the destination ID. The CPU processes all packets destined to the CLIP interface address. The system treats other packets with destination addresses associated with this network (but not to the interface address) as if they are from an unknown host.

A circuitless IP or CLIP address is a logical IP address for network management, as well as other purposes. The CLIP is typically a host address (with a 32 bit subnet mask). Configure the OSPF router ID to the configured CLIP address. By default, the BGP router ID is automatically equivalent to the OSPF router ID.

For information about how to configure CLIP interfaces, see [Configuring IPv4 Routing for VOSS](#).

BGP Configuration Considerations and Limitations

Use the information in this section to help you configure BGP on your switch, which supports BGPv4 as described in RFC 1771.

BGP Implementation Guidelines

The following list provides guidelines to successfully implement BGP:

- BGP does not operate with an IP router in nonforwarding (host-only) mode. Make sure that the routers you want BGP to operate with are in forwarding mode.
- If you use BGP for a multihomed AS (one that contains more than a single exit point), use OSPF for your IGP and BGP for your sole exterior gateway protocol, or use intra-AS iBGP routing.
- If OSPF is the IGP, use the default OSPF tag construction. Using EGP or modifying the OSPF tags makes network administration and proper configuration of BGP path attributes difficult.
- For routers that support both BGP and OSPF, the OSPF router ID and the BGP identifier must be the same IP address. The BGP router ID automatically uses the OSPF router ID.
- In configurations where BGP speakers reside on routers that have multiple network connections over multiple IP interfaces (the typical case for iBGP speakers), consider using the address of the circuitless (virtual) IP interface as the local peer address. In this configuration, you ensure that BGP is reachable as long as an active circuit exists on the router.
- By default, BGP speakers do not advertise or inject routes into the IGP. You must configure route policies to enable route advertisement.
- Coordinate routing policies among all BGP speakers within an AS so that every BGP border router within an AS constructs the same path attributes for an external path.
- Configure accept and announce policies on all iBGP connections to accept and propagate all routes. Make consistent routing policy decisions on external BGP connections.

Minimum Requirements

You must configure the following minimum parameters:

- router ID
- local AS number
- enable BGP globally
- BGP neighbor peer session: remote IP addresses
- enable BGP peers
- When you use both BGP and OSPF, the OSPF and BGP router ID must be the same.

The router ID must be a valid IP address of an IP interface on the router or a CLIP address. BGP update messages use this IP address. By default, the BGP router ID automatically uses the OSPF router ID.

You cannot configure the BGP router ID if you configure BGP before you configured the OSPF router ID. You must first disable BGP, configure the OSPF route ID, and then enable BGP globally.

You can add BGP policies to the BGP peer configuration to influence route decisions. BGP policies apply to the peer through the soft-reconfiguration commands.

After you configure the switch for BGP, some parameter changes can require you to enable or disable the BGP global state or the neighbor admin-state.

You can dynamically modify BGP policies. On the global level, the BGP redistribution command has an apply parameter that causes the policy to take effect after you issue the command.

BGP Neighbor Maximum Prefix Configuration

By default, the maximum prefix parameter limits 12 000 NLRI messages for each neighbor. The maximum prefix parameter limits the number of routes that the switch can accept.

The maximum prefix parameter prevents large numbers of BGP routes from flooding the network if you implement an incorrect configuration. You can assign a value to the maximum prefix limit, including 0 (0 means unlimited routes). When you configure the maximum prefix value, consider the maximum number of active routes that your equipment configuration can support.

BGP and OSPF Interaction

RFC1745 defines the interaction between BGP and OSPF when OSPF is the IGP within an autonomous system. For routers that use both protocols, the OSPF router ID and the BGP ID must be the same IP address. You must configure a BGP route policy to allow BGP advertisement of OSPF routes.

Interaction between BGPv4 and OSPF can advertise supernets to support CIDR. BGPv4 supports interdomain supernet advertisements; OSPF can carry supernet advertisements within a routing domain.

BGP and Internet Peering

By using BGP, you can perform Internet peering directly between the switch and another edge router. In such a scenario, you can use each switch for aggregation and link it with a Layer 3 edge router, as shown in the following figure.

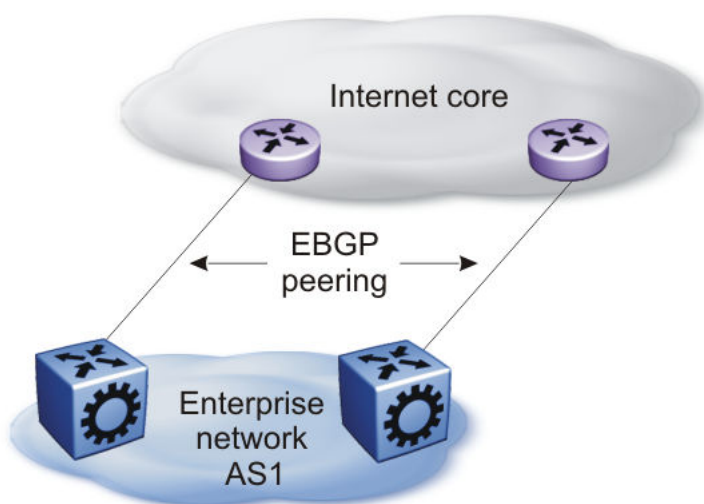


Figure 17: BGP and Internet peering

In cases where the Internet connection is single-homed, to reduce the size of the routing table, it is recommended that you advertise Internet routes as the default route to the IGP.

For route scaling information, see [Release Notes for XA1400 Series](#).

Routing Domain Interconnection with BGP

You can implement BGP so that autonomous routing domains, such as OSPF routing domains, connect. This connection allows the two different networks to begin communicating quickly over a common infrastructure, thus providing additional time to plan the IGP merger. Such a scenario is particularly effective when you need to merge two OSPF area 0.0.0.0s, as shown in the following figure.

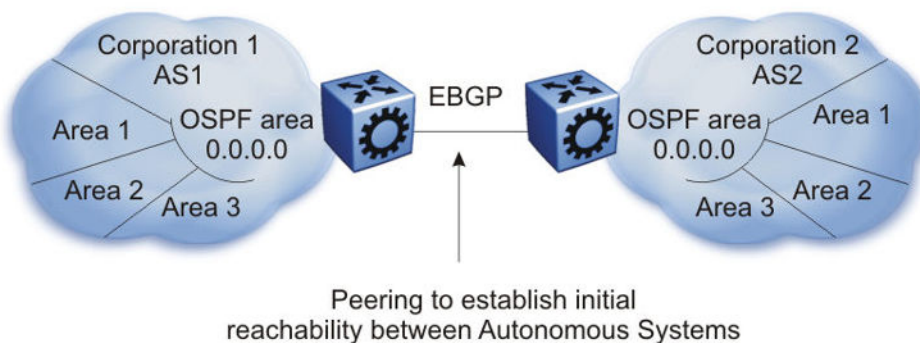


Figure 18: Routing Domain Interconnection with BGP

BGP and Edge Aggregation

You can perform edge aggregation with multiple point of presence or edge concentrations. The switch supports 12 pairs (peering services). You can use BGP to inject dynamic routes rather than using static routes or RIP (see the following figure).

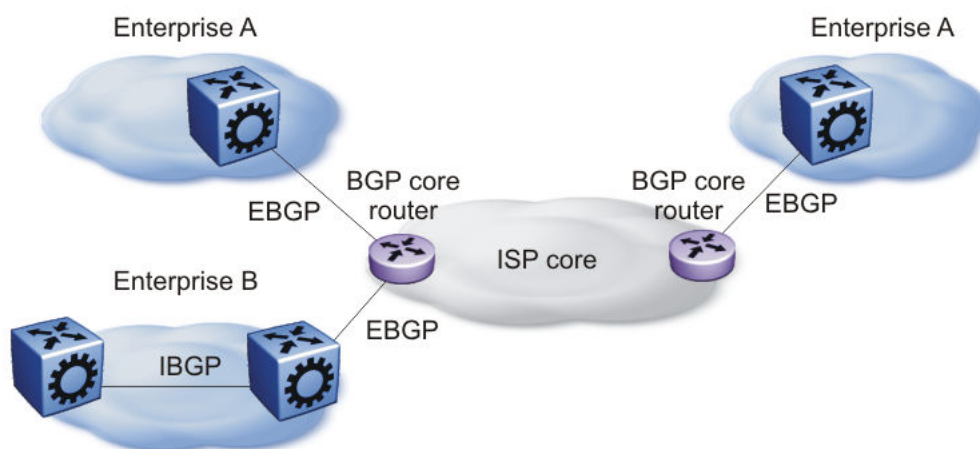


Figure 19: BGP and Edge Aggregation

BGP and ISP Segmentation

You can use the platform as a peering point between different regions or autonomous systems (AS) that belong to the same ISP. In such cases, you can define a region as an OSPF area, an AS, or a part of an AS.

You can divide the AS into multiple regions that each run different IGPs. Interconnect regions logically by using a full iBGP mesh. Each region then injects its IGP routes into iBGP and also injects a default route inside the region. For destinations that do not belong to the region, each region defaults to the BGP border router.

Use the community parameter to differentiate between regions. To provide Internet connectivity, this scenario requires you to make your Internet connections part of the central iBGP mesh (see the following figure).

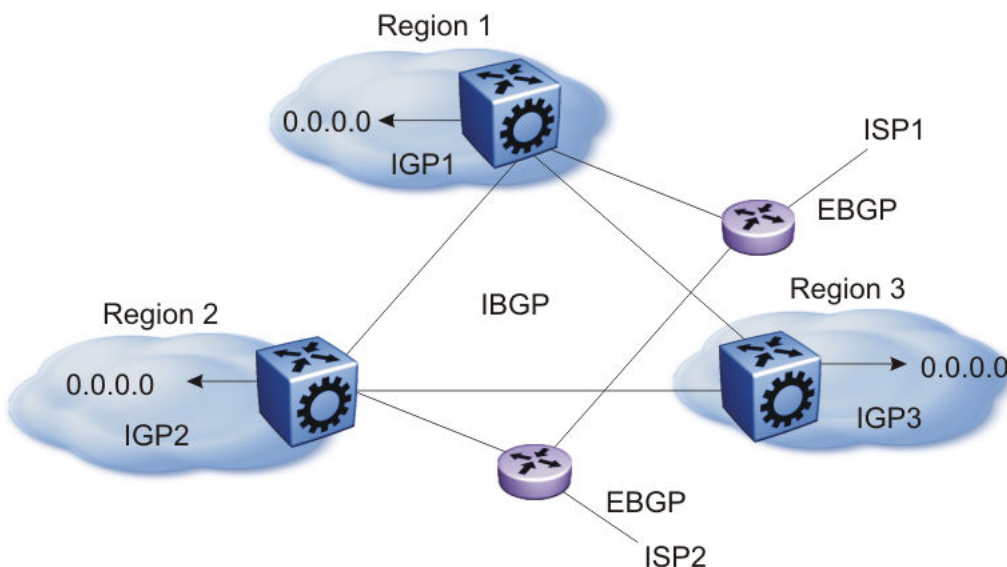


Figure 20: Multiple Regions Separated by iBGP

In the preceding figure, consider the following:

- The AS is divided into three regions that each run different and independent IGPs.
- Regions logically interconnect by using a full-mesh iBGP, which also provides Internet connectivity.
- Internal non-BGP routers in each region default to the BGP border router, which contains all routes.
- If the destination belongs to another region, the traffic is directed to that region; otherwise, the traffic is sent to the Internet connections according to BGP policies.

To configure multiple policies between regions, represent each region as a separate AS. Implement eBGP between autonomous systems, and implement iBGP within each AS. In such instances, each AS injects its IGP routes into BGP, where they are propagated to all other regions and the Internet.

The following figure shows the use of eBGP to join several autonomous systems.

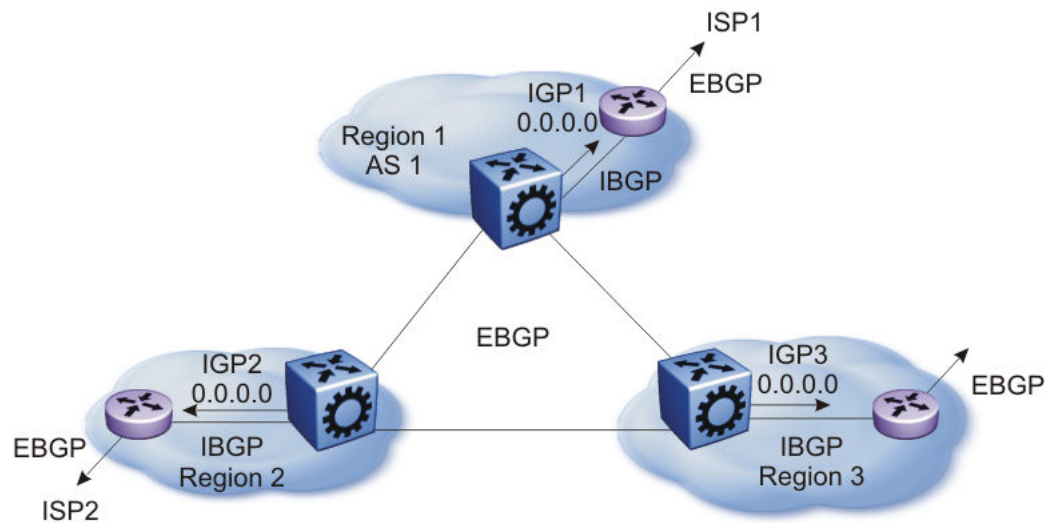


Figure 21: Multiple regions Separated by eBGP

You can obtain AS numbers from the Inter-Network Information Center (NIC) or use private AS numbers. If you use private AS numbers, be sure to design your Internet connectivity carefully. For example, you can introduce a central, well-known AS to provide interconnections between all private autonomous systems and the Internet. Before it propagates the BGP updates, this central AS strips the private AS numbers to prevent them from leaking to providers.

The following figure illustrates a design scenario in which you use multiple OSPF regions to enable peering with the Internet.

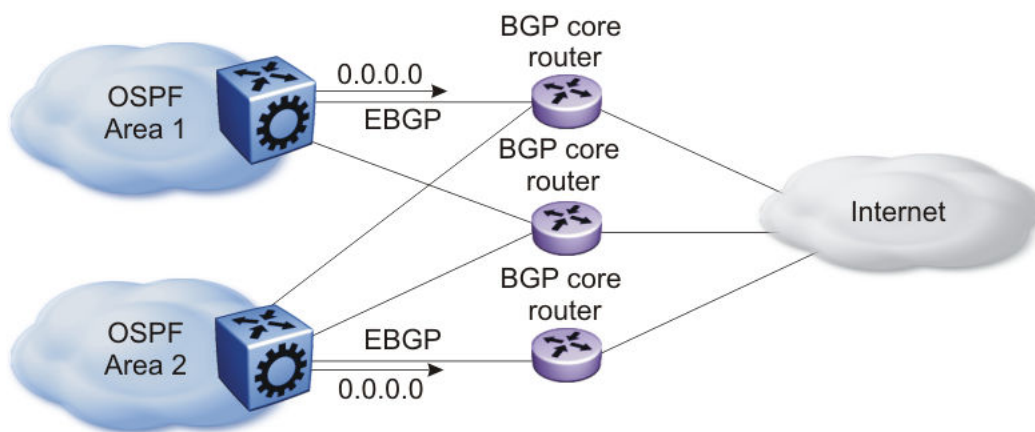


Figure 22: Multiple OSPF Regions Peering with the Internet

BGP Peers

The following list provides rules related to BGP peers:

- Only metric (=MED) attribute is applied to the output policy if its BGP peer is IBGP
- metric (=MED) and community attributes are applied to output policy if its BGP peer is EBGP

- To influence EBGp and IBGP peers with all applicable BGP attributes, configure **route-map** as an option to **neighbor** command, for example, `neighbor 192.0.2.2 out-route-map policy1`

BGP and Route Aggregation

When the attribute-map is configured with aggregate command, community and metric attributes are set, while the origin attribute is not set.

BGP Session Flapping when IPv6 Forwarding is Enabled or Disabled

In a BGP session that is established with IPv4 and IPv6 capability, disabling or enabling IPv6 forwarding results in BGP session flapping due to capability negotiation. The flapping session in turn affects the IPv4 routing through BGP and the BGP session gets terminated. Ultimately, a capability negotiation takes place to re-establish the IPv4 and IPv6 capable session.

Chapter 4: BGP configuration using CLI

Configure the Border Gateway Protocol (BGP) to create and maintain an interdomain routing system that guarantees loop-free routing information between autonomous systems (AS).

For information about how to configure route policies for BGP, see [Configuring IPv4 Routing for VOSS](#).

Configuring BGP Globally

Configure BGP globally to enable BGP on the switch and determine how BGP operates.

Before you begin

- To configure the suppress-map, advertise-map, or attribute-map options, the route policy for those options must exist.
- For initial BGP configuration, you must know the AS number.
- You configure BGP on a VRF instance the same way you configure the GlobalRouter, except that you must use VRF Router Configuration mode and the prefix `ip bgp`. The VRF must have an RP Trigger of BGP.

Note:

The following features are not currently supported on non-default VRFs:

- iBGP
- route refresh

Procedure

1. Enter Global Configuration mode:

```
enable
configure terminal
```

2. Specify the AS number and enable BGP:

```
router bgp [WORD <0-11>] [enable]
```

The AS number parameter only applies to VRF 0.

*** Note:**

This command applies only on VRF 0. To enable BGP globally on other VRFs, use the **ip bgp enable** command. You must configure BGP locally before you configure it globally.

3. Access Router BGP Configuration mode:

```
router bgp
```

4. Configure BGP variables or accept the default values.

Example

Specify the AS number and enable BGP:

```
Switch(config)#router bgp 3 enable
```

Access Router BGP Configuration mode:

```
Switch(config)#router bgp
```

```
Switch(router-bgp)#
```

Variable Definitions

Use the data in the following table to use the **router bgp** command.

Variable	Value
<i>WORD</i> <0-11>	Specifies the AS number. You cannot enable BGP until you change the local AS to a value other than 0.
enable	Enables BGP on the router.

Use the data in the following table to use the BGP variables in BGP Router Configuration mode.

Variable	Value
aggregate-address <i>WORD</i> <1-256>	Specifies an IP address and its length in the form {a.b.c.d/len}, or an IPv6 address and its length in the form {ipv6addr/len}.
auto-peer-restart enable	Enables the process that automatically restarts a connection to a BGP neighbor. The default value is enable.
auto-summary	When enabled, BGP summarizes networks based on class limits, for example, Class A, B, and C networks. The default value is enable.
bgp always-compare-med	Enables the comparison of the multi-exit discriminator (MED) parameter for paths from neighbors in different autonomous systems. The system prefers a path with a lower MED over a path with a higher MED. The default value is disable.
bgp client-to-client reflection	Enables or disables route reflection between two route reflector clients. This variable applies only if the route reflection value is

Table continues...

Variable	Value
	enable. The default value is disable. You can enable route reflection even when clients are fully meshed. This variable only applies to VRF 0. Example: Switch(router-bgp)# bgp client-to-client reflection System Response: Restart or soft-restart BGP for the change to take effect.
bgp cluster-id {A.B.C.D}	Configures a cluster ID. This variable applies only if the route reflection value is enable, and if multiple route reflectors are in a cluster. {A.B.C.D} is the IP address of the reflector router. This variable only applies to VRF 0. Example: Switch(router-bgp)# bgp cluster-id 0.0.0.0
bgp confederation identifier <0-4294967295> [peers WORD<0-255>]	Configures a BGP confederation. identifier<0-4294967295> specifies the confederation identifier. Use 0–65535 for 2-byte AS and <0-4294967295> for 4-byte AS. peers WORD<0-255> lists adjoining autonomous systems that are part of the confederation in the format (5500,65535,0,10,...,...). Use quotation marks (") around the list of autonomous systems.  Note: Use this command only on VRF 0. Example: Switch(router-bgp)# bgp confederation identifier 1 peers "20 30 40"
bgp default local-preference <0-2147483647>	Specifies the default value of the local preference attribute. The default value is 0. You must disable BGP before you can change the default value. Example: Switch(router-bgp)# bgp default local-preference 2-12
bgp deterministic-med enable	Enables deterministic MED. Example: Switch(router-bgp)# bgp deterministic-med enable
bgp multiple-paths <1-8>	Configures the maximum number of equal-cost-paths that are available to a BGP router by limiting the number of equal-cost-paths the routing table can store. The default value is 1. Example: Switch(router-bgp)# bgp multiple-paths 4

Table continues...

Variable	Value
	* Note: Configuring the bgp multiple-paths variable does not affect existing routes. The routing table does not show ECMP routes; instead only one route is shown in the routing table. To view Equal-Cost Multipath (ECMP) routes, receive the routes after executing the bgp multiple-paths variable, or toggle the BGP state. The number of equal-cost-paths supported can differ by hardware platform. For more information, see Release Notes for XA1400 Series.
comp-bestpath-med-confed enable	When enabled, compares MED attributes within a confederation. The default value is disable. This variable only applies to VRF 0. Example: Switch(router-bgp)# comp-bestpaht-med-confed enable Restart or soft-restart BGP for the change to take effect
debug-screen <off on>	Displays debug messages on the console, or saves them in a log file. Disable BGP screen logging (off) or enable BGP screen logging (on). Example: Switch(router-bgp)# debug-screen on System Response:BGP Screen Logging is On
default-information originate	Enables the advertisement of a default route to peers, if the route exists in the routing table. The default value is disable.
default-information ipv6-originate	Enables the advertisement of an IPv6 default route to peers, if the route exists in the routing table. The default value is disable.
default-metric <-1-2147483647>	Configures a value to send to a BGP neighbor to determine the cost of a route a neighbor uses. A default metric value solves the problems associated with redistributing routes that use incompatible metrics. For example, whenever metrics do not convert, using a default metric provides a reasonable substitute and redistribution proceeds. Use this option in conjunction with the redistribute commands so the current routing protocol uses the same metric for all redistributed routes. The default value is 0.
flap-dampening enable	Enables route suppression for routes that flap on and off. The default value is disable.
global-debug mask WORD<1-100>	Displays specified debug information for BGP global configurations. The default value is none. • <WORD 1-100> is a list of mask choices separated by commas with no space between choices.

Table continues...

Variable	Value
	Mask choices are: • none disables all debug messages. • all enables all debug messages. • error enables display of debug error messages. • packet enables display of debug packet messages. • event enables display of debug event messages. • trace enables display of debug trace messages. • warning enables display of debug warning messages. • state enables display of debug state transition messages. • init enables display of debug initialization messages. • filter enables display of debug messages related to filtering. • update enables display of debug messages related to sending and receiving updates. Example: Switch(router-bgp)# global-debug mask event, trace, warning, state
ibgp-report-import-rt enable	Configures BGP to advertise imported routes to an interior BGP (iBGP) peer. This variable enables or disables advertisement of nonBGP imported routes to other iBGP neighbors. The default value is enable.
ignore-illegal-rtrid enable	When enabled, BGP overlooks an illegal router ID. For example, you can configure this variable to enable or disable the acceptance of a connection from a peer that sends an open message using a router ID of 0 (zero). The default value is enable.
neighbor-debug-all mask WORD<1-100>	Displays specified debug information for BGP neighbors. The default value is none. For mask options, see the global-debug mask WORD<1-100> variable. Example: Switch(router-bgp)# neighbor-debug-all mask error, packet, event.trace, state, filter
no-med-path-is-worst enable	Enables BGP to treat an update without a MED attribute as the worst path. The default value is disable.
quick-start enable	Enables the quick-start flag for exponential backoff.
route-reflector enable	Enables the reflection of routes from iBGP neighbors. The default value is disable. This variable only applies to VRF 0.
route-refresh	Enables or disables route refresh. If enabled, a route refresh request received by a BGP speaker causes the speaker to resend all route updates it contains in its database that are eligible for the peer that issues the request.

Table continues...

Variable	Value
	This variable only applies to VRF 0.
router-id {A.B.C.D}	Specifies the BGP router ID in IP address format. This variable only applies to VRF 0.
synchronization	Enables the router to accept routes from BGP peers without waiting for an update from the IGP. The default value is enable.
traps enable	Enables BGP traps.

Job aid

Use debug command values to control debug messages for global BGP message types, and for message types associated with a specified BGP peer or peer group.

Tip:

The following tips can help you use the debug commands:

- Display debug commands for multiple mask choices by entering the mask choices separated by commas, with no space between choices.
- To end (disable) the display of debug messages, use the mask choice of none.
- You can save debug messages in a log file, or you can display the messages on your console using the debug-screen command.

For more information about the logged debug messages, see [Alarms and Logs Reference for VOSS](#).

Configuring 4-byte AS numbers

Configure Autonomous System (AS) numbers using the 4-byte format and represent the numbers in octets.

Before you begin

- You cannot modify the global BGP configuration unless BGP is disabled.
- Configure the local AS number at Global Router (VRF0) only.
- Make sure that you define AS numbers in policies the same way that you configure them for the router. The AS list for the route policies accepts AS number only in the **asplain** format. If you create policies using **asplain** and configure the switch with **asdot**, the match will not occur.

About this task

Use BGP 4-byte AS numbers to ensure the continuity of loop-free inter-domain routing information between autonomous systems and to control the flow of BGP updates as 2-byte AS numbers will

deplete soon. AS Plain notation format is the default and the preferred form of representing 4-byte AS numbers over the AS dot notation format.

You have an option to configure AS dot notation format as well. With AS dot notation, analyzing and troubleshooting any issues encountered becomes difficult as it is incompatible with the regular expressions used by most of the network providers.

Procedure

1. Enter Global Configuration mode:

```
enable
configure terminal
```

2. Disable BGP to change the AS number format.

```
no router bgp enable
```

3. Enable the 4-byte AS numbering format.

```
router bgp as-4-byte enable
```

4. To use the dotted octet notation, enable as-dot.

```
router bgp as-dot enable
```

5. Configure the 4-byte AS number and enable BGP. If you have enabled as-dot, enter the AS number in octet.

```
router bgp WORD<0-11> enable
```

6. Access Router BGP Configuration mode:

```
router bgp
```

7. **(Optional)** Configure BGP confederation identifier.

```
bgp confederation identifier <0-4294967295>
```

8. **(Optional)** Configure BGP confederation peers.

```
bgp confederation peers WORD<0-255>
```

Example

Disable BGP to change the AS number format.

```
Switch(config)# no router bgp enable
```

Enable the 4-byte AS numbering format.

```
Switch(config)# router bgp as-4-byte enable
```

To use the dotted octet notation, enable as-dot.

```
Switch(config)# router bgp as-dot enable
```

Configure the 4-byte AS number and enable BGP.

```
Switch(config)# router bgp 65536 enable
```

Variable Definitions

Use the data in the following table to use the **router bgp** command.

Variable	Value
as-4-byte <enable>	Enables the switch for using 4 byte numbers for an autonomous system (AS). The default value is disable.
as-dot <enable>	Enables or disables representing AS numbers in octets. The default is disable so the switch uses the plain notation format. If you enable the 4-byte-as and as-dot parameters, enter numbers in the range of 1.0 to 65535.65535. The default value is disable.  Note: This parameter is not supported with BGP+.
WORD <0–11> enable	Sets the local autonomous system (AS) number. You cannot change local-as when BGP is set to enable. - To set a 2-byte local AS number, enter a local-as number in the range of 0 to 65535. - To set a 4-byte local-as number, enable the 4-byte as variable and enter a number in the range of 0 to 4294967295.  Note: If as-4-byte is set to false, the range for AS number is 0–65535 and if as-4-byte is set to true, the range is 0–4294967295. If you enable as-dot, enter the AS number in octets in the range of 1.0 to 65535.65535.  Note: This parameter is not supported with BGP+.

Configuring Aggregate Routes

Configure aggregate routes so that the router advertises a single route (aggregate route) that represents all destinations. Aggregate routes also reduce the size of routing tables.

Before you begin

- Disable BGP before you enable aggregation.
- You need the appropriate aggregate address and mask.
- If required, policies exist.

- You configure BGP on a VRF instance the same way you configure the GlobalRouter, except that you must use VRF Router Configuration mode and the prefix `ip bgp`. The VRF must have an RP Trigger of BGP.

 Note:

The following features are not currently supported on non-default VRFs:

- iBGP
- route refresh

Procedure

1. Enter BGP Router Configuration mode:

```
enable
configure terminal
router bgp
```

2. Enable BGP aggregation:

```
bgp aggregation enable
```

3. Add an aggregate route to the routing table:

```
aggregate-address WORD<1-256> {advertise-map WORD<0-1536>} [as-set]
[attribute-map WORD<0-1536>] [summary-only] [suppress-map WORD<0-1536>]
```

4. Exit to Global Configuration mode:

```
exit
```

5. Enable BGP:

```
router bgp [<0-65535>] [enable]
```

Example

Add an aggregate route to the routing table:

```
Switch(router-bgp)# aggregate-address 2001:DB8::/32 advertise-map map1
attribute-map map2
```

Enable BGP:

```
Switch(router-bgp)# router bgp 4 enable
```

Variable definitions

Use the data in the following table to use the `aggregate-address` command.

Variable	Value
advertise-map WORD<0-1536>	Specifies the route map name for route advertisements.
as-set	Enables autonomous system information. The default value is disable.
attribute-map WORD<0-1536>	Specifies the route map name.
WORD <1–256>	Specifies an IP address and its length in the appropriate form. The value must be entered in the format a.b.c.d/len or ipv6addr/len.
summary-only	Enables the summarization of routes not included in routing updates. This variable creates the aggregate route and suppresses advertisements of more specific routes to all neighbors. The default value is disable.
suppress-map WORD<0-1536>	Specifies the route map name for the suppressed route list.

Use the data in the following table to use the **router bgp** command.

Variable	Value
<0-65535>	Specifies the AS number. You cannot enable BGP until you change the local AS to a value other than 0.
enable	Enables BGP on the router.

Configuring allowed networks

Configure network addresses to determine the network addresses that BGP advertises. The allowed addresses determine the BGP networks that originate from the switch.

Before you begin

- You configure BGP on a VRF instance the same way you configure the GlobalRouter, except that you must use VRF Router Configuration mode and the prefix **ip bgp**. The VRF must have an RP Trigger of BGP. Not all parameters are configurable on non-default VRFs.

Procedure

- Enter BGP Router Configuration mode:

```
enable
configure terminal
router bgp
```

- Specify IGP network prefixes for BGP to advertise:

```
network <WORD 1-256> [metric <0-65535>]
```

Example

Specify IGP network prefixes for BGP to advertise:

```
Switch(router-bgp)# network 2001:DB8::/32 metric 32
```

Variable definitions

Use the data in the following table to use the **network** command.

Variable	Value
WORD <1–256>	Specifies an IP address and its length in the appropriate form.
metric <0–65535>	Specifies the metric to use when the system sends an update for the routes in the network table. The metric configures the MED for the routes advertised to eBGP peers. The range is 0–65535.

Configuring BGP peers or peer groups

Configure peers and peer groups to simplify BGP configuration and makes updates more efficient.

BGP speakers can have many neighbors configured with similar update policies. For example, many neighbors use the same distribute lists, filter lists, outbound route maps, and update source. Group the neighbors that use the same update policies into peer groups and peer associations.

Before you begin

- If required, route policies exist.
- You configure BGPv4 on a VRF instance the same way you configure the GlobalRouter, except that you must use VRF Router Configuration mode and the prefix **ip bgp**. The VRF must have an RP Trigger of BGP. **iBGP** and **route refresh** are not currently supported on non-default VRFs.

About this task

Many of the command variables in this procedure use default values. You can accept the default values or change them to customize the configuration.

Procedure

1. Enter BGP Router Configuration mode:

```
enable
configure terminal
router bgp
```

2. Create a peer or peer group:

```
neighbor WORD<0-1536>
```

3. Apply a route policy to all incoming routes:

For BGPv4: `neighbor WORD<0-1536> in-route-map WORD<0-256>`

For BGPv6: `neighbor WORD<0-1536> ipv6-in-route-map WORD<0-256>`

4. Apply a route policy to all outgoing routes:

For BGPv4: `neighbor WORD<0-1536> out-route-map WORD<0-256>`

For BGPv6: `neighbor WORD<0-1536> ipv6-out-route-map WORD<0-256>`

5. **(Optional)** Configure the source IP address:

```
neighbor WORD<0-1536> update-source WORD<1-256>
```

6. Enable MD5 authentication (for BGPv4):

```
neighbor WORD<0-1536> MD5-authentication enable
```

7. Specify an MD5 authentication password (for BGPv4):

```
neighbor password <nbr_ipaddr|peer-group-name> WORD<0-1536>
```

8. Change the default values for other command variables as required.

9. Enable the configuration:

```
neighbor WORD<0-1536> enable
```

Example

Create a peer or a peer group:

```
Switch(router-bgp)# neighbor peergroupa
```

Apply a route policy (in-route-map or out-route-map) to all incoming or outgoing routes:

```
Switch(router-bgp)# neighbor peergroupa in-route-map map1 out-route-map  
map2
```

Configure the source IP address:

```
Switch(router-bgp)# neighbor peergroupa update-source 192.0.2.1
```

Enable MD5 authentication:

```
Switch(router-bgp)# neighbor peergroupa MD5-authentication enable
```

Specify an MD5 authentication password:

```
Switch(router-bgp)# neighbor password peergroupa password
```

Enable the configuration:

```
Switch(router-bgp)# neighbor peergroupa enable
```

Variable definitions

Use the data in the following table to use the **neighbor** command.

Variable	Value
address-family <ipv6>	Enables the IPv6 address family on BGP neighbor. Switch(router-bgp)# neighbor peergroupa address-family ipv6
advertisement-interval <5-120>	Specifies the time interval, in seconds, that transpires between each transmission of an advertisement from a BGP neighbor. The default value is 5 seconds. Switch(router-bgp)# neighbor peergroupa advertisement-interval 26 enable The route advertisement interval feature is implemented using the time stamp that indicates when each route is advertised. The time stamp is marked to each route so that the route advertisement interval is compared to the time stamp and BGP is then able to make a decision about whether the route advertisement can be sent or it should be delayed when a better route is received. This feature does not work for a withdraw route because the route entry is already removed when the processing route advertisement is sent and the time stamp marked in the route entry cannot be obtained.
default-ipv6-originate	Enables IPv6 BGP neighbor default originate. Switch(router-bgp)# neighbor peergroupa default-ipv6-originate
default-originate	Enables the switch to send a default route advertisement to the specified neighbor. A default route does not need to be in the routing table. The default value is disable. Do not use this command if default-information originate is globally enabled. Switch(router-bgp)# neighbor peergroupa default-originate enable peer-group test
ebgp-multihop	Enables a connection to a BGP peer that is more than one hop away from the local router. The default value is disable. Switch(router-bgp)# neighbor peergroupa ebgp-multihop retry-interval 3 timers 4 5
enable	Enables the BGP neighbor.
in-route-map WORD<0-256>	Applies a route policy rule to all incoming routes that are learned from, or sent to, the peers or peer groups of the local router. The local BGP router is the BGP router that allows or disallows routes and configures attributes in incoming or outgoing updates.

Table continues...

Variable	Value
	WORD<0-256> is an alphanumeric string length (0–256 characters) that indicates the name of the route map or policy. <pre>Switch(router-bgp)# neighbor peergroupa in- route-map map1 address-family ipv6</pre>
ipv6-in-route-map WORD <0–256>	Creates IPv6 in route map. WORD <0–256> specifies the route map name in the range of 0 to 256 characters. <pre>Switch(router-bgp)# neighbor peergroupa ipv6- in-route-map map1</pre>
ipv6-max-prefix <0-2147483647>	Configures a limit on the number of routes that the router can accept from a neighbor. The default value is 12000 routes. A value of 0 (zero) indicates that no limit exists.
ipv6-out-route-map WORD <0–256>	Creates IPv6 out route map. WORD <0–256> specifies the route map name in the range of 0 to 256 characters. <pre>Switch(router-bgp)# neighbor peergroupa ipv6- out-route-map map2</pre>
max-prefix <0-2147483647>	Configures a limit on the number of routes that the router can accept from a neighbor. The default value is 12000 routes. A value of 0 (zero) indicates that no limit exists. <pre>Switch(router-bgp)# neighbor peergroupa max- prefix 158 in-route-map map1 out-route-map map2</pre>
MD5-authentication enable	Enables TCP MD5 authentication between two peers. The default value is disable.
neighbor-debug mask WORD<1-100>	Displays specified debug information for a BGP peer. The default value is none. <WORD 1-100> is a list of mask choices separated by commas with no space between choices. For example: {<mask>,<mask>,<mask>...}. Mask choices are: • none disables all debug messages. • all enables all debug messages. • error enables display of debug error messages. • packet enables display of debug packet messages. • event enables display of debug event messages. • trace enables display of debug trace messages. • warning enables display of debug warning messages. • state enables display of debug state transition messages. • init enables display of debug initialization messages. • filter enables display of debug messages related to filtering.

Table continues...

Variable	Value
	• update enables display of debug messages related to sending and receiving updates. <pre>Switch(router-bgp)# neighbor peergroupa neighbor-debug-mask event,trace,warning,state</pre>
next-hop-self	When enabled, specifies that the next-hop attribute in an iBGP update is the address of the local router or the router that generates the iBGP update. The default value is disable. You can only configure this variable if the neighbor is disabled. <pre>Switch(router-bgp)# neighbor peergroupa next-hop-self out-route-map map2 peer-group peergroupb</pre>
out-route-map WORD<0-256>	Applies a route policy rule to all outgoing routes that are learned from, or sent to, the peers or peer groups of the local router. The local BGP router is the BGP router that allows or disallows routes and configures attributes in incoming or outgoing updates. WORD<0-256> is an alphanumeric string length (0–256 characters) that indicates the name of the route map or policy.
peer-group <WORD 0-1536>	Adds a BGP peer to the specified subscriber group. You must create the specified subscriber group before you use this command.
remote-as <WORD 0-11>	Configures the remote AS number of a BGP peer or a peer-group. You must disable the admin-state before you can configure this variable. <pre>Switch(router-bgp)# neighbor peergroupa remote-as As-number</pre> <WORD 0-11> is an alphanumeric string length (0–11 characters) that indicates the AS number.
remove-private-as enable	Strips private AS numbers when an update is sent. The default value is enable.
retry-interval <1-65535>	Configures the time interval, in seconds, for the ConnectRetry timer. The default value is 120 seconds. <pre>Switch(router-bgp)# neighbor 198.51.100.2 retry-interval 34</pre> You can configure the retry interval for BGP neighbors only; you cannot configure the retry interval for BGP peer groups.
route-reflector-client	Configures the specified neighbor or group of neighbors as a route reflector client. The default value is disable. All configured neighbors become members of the client group and the remaining iBGP peers become members of the nonclient group for the local route reflector.

Table continues...

Variable	Value
	 Note: This variable only applies to VRF 0. <pre>Switch(router-bgp)# neighbor</pre>
route-refresh	Enables route refresh for the BGP peer. If enabled, a route refresh request received by a BGP speaker causes the speaker to resend all route updates it contains in its database that are eligible for the peer that issues the request.  Note: This variable only applies to VRF 0.
send-community	Enables the switch to send the update message community attribute to the specified peer. The default value is disable.
soft-reconfiguration-in enable	Enables the router to relearn routes from the specified neighbor or group of neighbors without restarting the connection after the policy changes in the inbound direction. The default value is disable.
timers <0-21845> <0-65535>	Configures timers, in seconds, for the BGP speaker for this peer. <0-21845> is the keepalive time. The default is 60. It is recommended that you configure a value of 30 seconds. <0-65535> is the hold time. The default is 180. <pre>Switch(router-bgp)# neighbor peergroupa timers 4 6</pre>
update-source WORD<1–256>	Specifies the source IPv4 address {A.B.C.D.} or IPv6 address to use when the system sends BGP packets to this peer or peer group. You must disable the admin-state before you can configure this variable. <pre>Switch(router-bgp)# neighbor peergroupa update-source 192.0.2.2 weight 560</pre>
weight <0-65535>	Specifies the weight of a BGP peer or peer group, or the priority of updates the router can receive from that BGP peer. The default value is 0. If you have particular neighbors that you want to use for most of your traffic, you can assign a higher weight to all routes learned from that neighbor.
WORD<0-1536>	Specifies the peer IP address or the peer group name.

Configuring a BGP peer or peer group password

Use this procedure to configure a BGP peer or peer group password for Transmission Control Protocol (TCP) MD5 authentication between two peers.

Procedure

1. Enter BGP Router Configuration mode:

```
enable
configure terminal
router bgp
```

2. Assign a BGP peer or peer group password:

```
neighbor password <nbr_ipaddr|peer-group=name> WORD <0-1536>
```

Example

Assign a BGP peer or peer group password:

```
Switch(router-bgp)# neighbor password peergroupa password1
```

Variable definitions

Use the data in the following table to use **neighbor password <nbr_ipaddr|peer-group-name>** command.

Variable	Value
password <nbr_ipaddr peer-group-name> WORD <0-1536>	Specifies a password for TCP MD5 authentication between two peers. WORD <0-1536> is an alphanumeric string length from 0 to 1536 characters. To disable this option, use no operator with the command. To configure this option to the default value, use default operator with the command.

Configuring redistribution to BGP

Configure a redistribution entry to announce routes of a certain source protocol type into the BGP domain such as: DvR routes, static routes, Routing Information Protocol (RIP) routes, or direct routes. Use a route policy to control the redistribution of routes.

 Note:

When a route map with attributes set to `origin` and `local-pref` is applied to the BGP redistribute command, the attributes are not applied to the redistributed routes.

Before you begin

- If required, a route policy exists.

- You configure BGP on a VRF instance the same way you configure the GlobalRouter, except that you must use VRF Router Configuration mode and the prefix `ip bgp`. The VRF must have an RP Trigger of BGP. `iBGP` and `route refresh` are not currently supported on non-default VRFs.
- Before you redistribute DvR host routes to BGP, you must disable BGP aggregation and BGP auto-summarization of networks, using the commands `no ip bgp aggregation enable` and `no ip bgp auto-summary` respectively.

Disabling these settings ensures that all the DvR host routes are correctly advertised into BGP and are not summarized.

*** Note:**

When applying a route map to an inter-vrf redistribution, the route map and any associated IP prefix lists must be configured first on the source VRF before configuring the redistribute policy on the destination VRF.

Inter-vrf redistribution is not supported on IPv6 routes.

Procedure

1. Enter BGP Router Configuration mode:

```
enable
configure terminal
router bgp
```

2. Create a redistribution instance:

```
redistribute <direct|dvr|ipv6-direct|ipv6-isis|ipv6-static|isis|
ospf|ospfv3|rip|ripng|static>
```

*** Note:**

Redistribution of ripng routes into BGP is supported only on VRF 0.

3. If required, specify a route policy to govern redistribution:

```
redistribute <direct|dvr|ipv6-direct|ipv6-isis|ipv6-static|isis|
ospf|ospfv3|rip|ripng|static> route-map WORD<0-64> [vrf-src
WORD<1-16>]
```

4. If required, configure the route metric:

```
redistribute <direct|dvr|ipv6-direct|ipv6-isis|ipv6-static|isis|
ospf|ospfv3|rip|ripng|static> metric <0-65535> [vrf-src WORD<1-16>]
```

5. If required, configure the route metric-type:

```
redistribute <direct|dvr|ipv6-direct|ipv6-isis|ipv6-static|isis|
ospf|ospfv3|rip|ripng|static> metric-type live-metric [vrf-src
WORD<1-16>]
```

6. Enable the instance:

```
redistribute <direct|dvr|ipv6-direct|ipv6-isis|ipv6-static|isis|
ospf|ospfv3|rip|ripng|static> enable [vrf-src WORD<1-16>]
```

7. Exit BGP Router Configuration mode:

```
exit
```

8. Apply the redistribution instance configuration:

For IPv4: `ip bgp apply redistribute <direct|dvr|isis|ospf|rip|static> [vrf WORD<1-16>] [vrf-src <WORD 1-16>]`

For IPv6: `ipv6 bgp apply redistribute <direct|dvr|isis|ospf|rip|static> [vrf <WORD 1-16>]`

9. Apply BGP redistribution to a specific VRF:

```
ip bgp apply redistribute vrf WORD<1-16>
```

Changes do not take effect until you apply them.

10. View all routes (including DvR host routes) that are redistributed into BGP:

View routes redistributed from GRT to BGP:

For IPv4: `show ip bgp imported-routes`

For IPv6: `show bgp ipv6 imported-routes`

View routes redistributed to BGP for a specific VRF instance:

For IPv4: `show ip bgp imported-routes [vrf WORD<1-64>] [vrfids WORD<0-512>]`

For IPv6: `show bgp ipv6 imported-routes [WORD<1-256>] [vrf WORD<1-16>] [vrfids WORD<0-255>]`

Example

Example 1:

Redistribute direct routes from the VRF instance `source1` into BGP, in the GRT context.

Create a redistribution instance:

```
Switch:1>enable
Switch:1#configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch:1(config)#router bgp
Switch(router-bgp)#redistribute direct vrf-src source1
```

If required, specify a route policy to govern redistribution:

```
Switch(router-bgp)# redistribute direct route-map policy1 vrf-src source1
```

If required, configure the route metric:

```
Switch:1(router-bgp)# redistribute direct metric 4 vrf-src source1
```

Enable the instance:

```
Switch:1(router-bgp)# redistribute direct enable vrf-src source1
```

Exit BGP Router Configuration mode:

```
Switch:1(router-bgp)# exit
```

Apply the redistribution instance configuration:

```
Switch:1(config)# ip bgp apply redistribute direct vrf-src source1
```

Example 2:

Redistribute DvR routes from the GRT to BGP:

```
Switch:1>enable
Switch:1#configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch:1(config)#router bgpSwitch:1(router-bgp)#redistribute dvr
Switch:1(router-bgp)#redistribute dvr enable
Switch:1(router-bgp)#exit
Switch:1(config)#ip bgp apply redistribute dvr
```

View the host routes (including DvR host routes) that are redistributed from the GRT to BGP:

```
Switch:1(config)#show ip bgp imported-routes vrf vrf1
```

```
=====
                        BGP Imported Routes - VRF vrf1
=====
```

ROUTE	METRIC	COMMUNITY	LOCALPREF	NEXTHOP	ORIGIN
192.0.2.1/255.255.255.0	0	0	100	198.51.100.1	INC
192.0.2.2/255.255.255.0	0	0	100	198.51.100.1	INC
192.0.2.3/255.255.255.0	0	0	100	198.51.100.1	INC
...					
...					
...					

3 out of 763 Total Num of imported routes displayed

Example 3:

Redistribute DvR routes to BGP for the specific VRF instance vrf1:

```
Switch:1>enable
Switch:1#configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch:1(config)#router vrf vrf1
Switch:1(router-vrf)#ip bgp redistribute dvr
Switch:1(router-vrf)#ip bgp redistribute dvr enable
Switch:1(router-vrf)#exit
Switch:1(config)#ip bgp apply redistribute dvr vrf vrf1
```

View the DvR host routes that are redistributed to BGP for vrf vrf1:

```
Switch:1(config)#show ip bgp imported-routes vrf vrf1
```

```
=====
                        BGP Imported Routes - VRF vrf1
=====
```

ROUTE	METRIC	COMMUNITY	LOCALPREF	NEXTHOP	ORIGIN
192.0.2.4/255.255.255.0	0	0	100	203.0.113.1	INC
192.0.2.5/255.255.255.0	0	0	100	203.0.113.1	INC
192.0.2.6/255.255.255.0	0	0	100	203.0.113.1	INC
192.0.2.7/255.255.255.0	0	0	100	203.0.113.1	INC
192.0.2.8/255.255.255.0	0	0	100	203.0.113.1	INC
...					
...					

```
...
5 out of 675 Total Num of imported routes displayed
```

Example 4:

This example demonstrates redistribution of inter-VRF routes (both direct and DvR routes) to BGP, with a route policy configured.

Redistribute inter-VRF DvR routes between VRFs (with VRF IDs 10 and 30), to BGP.

```
Switch:1>enable
Switch:1#configure terminal
Switch:1(config)#router vrf 10
Switch:1(router-vrf)#ip prefix-list "test10" 192.0.2.0/24 ge 25 le 32
Switch:1(router-vrf)#route-map "test10" 1
Switch:1(router-vrf)#permit
Switch:1(router-vrf)#enable
Switch:1(router-vrf)#match network "test10"
Switch:1(router-vrf)#set metric 99
Switch:1(router-vrf)#exit

Switch:1(config)#router vrf 30
Switch:1(router-vrf)#ip bgp redistribute direct vrf-src 10
Switch:1(router-vrf)#ip bgp redistribute direct enable vrf-src 10
Switch:1(router-vrf)#ip bgp redistribute dvr vrf-src 10
Switch:1(router-vrf)#ip bgp redistribute dvr route-map "test10" vrf-src 10
Switch:1(router-vrf)#ip bgp redistribute dvr enable vrf-src 10
Switch:1(router-vrf)#exit

Switch:1(config)#ip bgp apply redistribute direct vrf 30 vrf-src 10
Switch:1(config)#ip bgp apply redistribute dvr vrf 30 vrf-src 10
```

Variable definitions

Use the data in the following table to use the **redistribute** and **ip bgp apply redistribute** commands.

Variable	Value
<direct dvr ipv6-direct ipv6-isis ipv6-static isis ospf ospfv3 rip ripng static >	Specifies the type of routes to redistribute (the protocol source).
enable	Enables the BGP route redistribution instance.
metric <0-65535>	Configures the metric to apply to redistributed routes.
metric-type live-metric	Configures the metric type to apply to redistributed routes. When you enable the live-metric option, when BGP redistributes static, RIP, OSPF, IS-IS, or DvR routes, the metric value is taken from the routing table and is set to the Path attributes as a MED value. By default, this option is disabled, which means the BGP MED value is not derived from the metric in the routing table.

Table continues...

Variable	Value
route-map WORD<0-64>	Configures the route policy to apply to redistributed routes.
vrf WORD<1-16>	Specifies the name of a VRF instance.
vrf-src WORD<1-16>	Specifies the source VRF instance by name for route redistribution.

Configuring redistribution to BGP+ for VRF 0

Configure an IPv6 redistribute entry to announce IPv6 routes of a certain source protocol type into the BGP domain, for example, static, OSPF, IS-IS, RIPng, or direct routes. Use a route policy to control the redistribution of routes.

* Note:

When a route map with attributes set to `origin` and `local-pref` is applied to the BGP redistribute command, the attributes are not applied to the redistributed routes.

Before you begin

- If required, a route policy exists.

Procedure

1. Enter BGP Router Configuration mode:

```
enable
configure terminal
router bgp
```

2. Create a redistribution instance:

```
redistribute <<irect|dvr|ipv6-direct|ipv6-isis|ipv6-static|isis|
ospf|ospfv3|rip|ripng|static>
```

3. If required, specify a route policy to govern redistribution:

```
redistribute <direct|dvr|ipv6-direct|ipv6-isis|ipv6-static|isis|
ospf|ospfv3|rip|ripng|static> route-map WORD <0-64>
```

4. If required, configure a route metric:

```
redistribute <direct|dvr|ipv6-direct|ipv6-isis|ipv6-static|isis|
ospf|ospfv3|rip|ripng|static> metric <0-65535>
```

5. Enable the instance:

```
redistribute <direct|dvr|ipv6-direct|ipv6-isis|ipv6-static|isis|
ospf|ospfv3|rip|ripng|static> enable
```

Unlike IPv4 redistribution, you do not need to manually apply the IPv6 redistribution instance. Once you enable the IPv6 redistribution instance, it is automatically applied.

Example

Specify a route policy to govern redistribution by using the following command:

```
Switch:1(router-bgp)#redistribute ipv6-direct route-map policy2
```

Variable definitions

Use the information in the following table to help you use the **redistribute** <ipv6-direct|ipv6-static|ospfv3|ipv6-isis|ripng>command.

Table 8: Variable definitions

Variable	Value
enable	Enables the BGP route redistribution instance. The default value is none. To configure this option to the default value, use default operator with the command. To disable this option, use no operator with the command.
metric<0–65535>	Configures the metric to apply to redistributed routes. The default value is 0. To configure this option to the default value, use default operator with the command.
route-map <Word 0-64>	Configures the route policy to apply to redistributed routes. The default value is none. To configure this option to the default value, use default operator with the command.

Job Aid

Use the data in the following table to know how route policies are used for BGP from IPv6 perspective.

Table 9: BGP for IPv6 Route Policy Support

	REDISTRIBUTE					ACCEPT	ANNOUNCE
	IPv6 Direct	IPv6 Static	OSPFv3	IPv6 IS-IS	RIPng	BGP	BGP
MATCH							

Table continues...

	REDISTRIBUTE					ACCEPT	ANNOUNCE
	IPv6 Direct	IPv6 Static	OSPFv3	IPv6 IS-IS	RIPng	BGP	BGP
as-path						Yes	Yes
community	Yes	Yes	Yes	Yes	Yes	Yes	Yes
community-exact						Yes	Yes
extcommunity						Yes	Yes
interface							
local-preference							
metric				Yes	Yes		
network				Yes	Yes		
next-hop				Yes	Yes		
protocol							
route-source						Yes	
route-type			Yes				Yes
tag							
vrf							
vrfids							
SET							
as-path						Yes	Yes
as-path-mode						Yes	Yes
automatic-tag							
community						Yes	Yes
community-mode						Yes	Yes
injectlist	Yes	Yes	Yes	Yes	Yes		
ip-preference							
local-preference						Yes	Yes
mask							
metric	Yes	Yes	Yes	Yes	Yes	Yes	Yes
metric-type							
metric-type-internal							
next-hop						Yes	Yes
nssa-pbit							
origin							Yes
origin-egp-as							

Table continues...

	REDISTRIBUTE					ACCEPT	ANNOUNC E
	IPv6 Direct	IPv6 Static	OSPFv3	IPv6 IS-IS	RIPng	BGP	BGP
tag							
weight						Yes	

Configuring AS path lists

Configure an AS path list to restrict the routing information a router learns or advertises to and from a neighbor. The AS path list acts as a filter that matches AS paths.

Before you begin

- You configure BGP on a VRF instance the same way you configure the GlobalRouter, except that you must use VRF Router Configuration mode and the prefix **ip bgp**. The VRF must have an RP Trigger of BGP. **iBGP** and **route refresh** are not currently supported on non-default VRFs.

Procedure

- Enter BGP Router Configuration mode:

```
enable
configure terminal
router bgp
```

- Create the path list:

```
ip as-list <1-1024> memberid <0-65535> <permit|deny> as-path WORD<0-1536>
```

Use this command for each member by specifying different member IDs.

Example

Create the path list:

```
Switch(config)# ip as-list 234 memberid 3456 permit as-path "5"
```

Variable definitions

Use the data in the following table to use the **ip as-list** command.

Variable	Value
<0-65535>	Specifies an integer value between 0–65535 that represents the regular expression entry in the AS path list.
<1-1024>	Specifies an integer value from 1–1024 that represents the AS-path list ID you want to create or modify.
<permit deny>	Permits or denies access for matching conditions.
WORD<0–1536>	Specifies the AS number as an integer value between 0–1536. Place multiple AS numbers within quotation marks (").

Configuring community lists

Configure community lists to specify permitted routes by using their BGP community. This list acts as a filter that matches communities or AS numbers.

Before you begin

- You configure BGP on a VRF instance the same way you configure the GlobalRouter, except that you must use VRF Router Configuration mode and the prefix `ip bgp`. The VRF must have an RP Trigger of BGP. `iBGP` and `route refresh` are not currently supported on non-default VRFs.

Procedure

1. Enter BGP Router Configuration mode:

```
enable
configure terminal
router bgp
```

2. Create a community list:

```
ip community-list <1-1024> memberid <0-65535> <permit|deny>
community-string WORD<0-256>
```

Example

Create a community list:

```
Switch(config)# ip community-list 1 memberid 4551 permit community-string
internet
```

Variable Definitions

Use the data in the following table to use the `ip community-list` command.

Variable	Value
<0-65535>	Specifies an integer value from 0–65535 that represents the member ID in the community list.
<1-1024>	Specifies an integer value from 1–1024 that represents the community list ID.
<permit deny>	Configures the access mode, which permits or denies access for matching conditions.
WORD<0-256>	Specifies the community as an alphanumeric string value with a string length from 0–256 characters. Enter this value in one of the following formats: • (AS num:community-value) • (well-known community string) Well known communities include: internet, no-export, no-advertise, local-as (known as NO_EXPORT_SUBCONFED).

Configuring extended community lists

Configure community lists to specify permitted routes by BGP extended community attributes, including route targets and sites of origin (SOO). This list acts as a filter that matches route targets and SOO.

Before you begin

- You configure BGP on a VRF instance the same way you configure the GlobalRouter, except that you must use VRF Router Configuration mode and the prefix `ip bgp`. The VRF must have an RP Trigger of BGP. `iBGP` and `route refresh` are not currently supported on non-default VRFs.

Procedure

1. Enter BGP Router Configuration mode:

```
enable
configure terminal
router bgp
```

2. Create an extended community list based on the route target attribute:

```
ip extcommunity-list <1-1024> memberId <0-65535> rt {<0-65535>
<0-2147483647>|<A.B.C.D> <0-65535>} [soo {<0-65535> <0-2147483647>|
<A.B.C.D> <0-65535>}]
```

You can optionally configure the SOO attributes at the end of the same command or you can configure the SOO separately using the syntax in the following step.

3. Create an extended community list based on the SOO attribute:

```
ip extcommunity-list <1-1024> memberId <0-65535> soo {<0-65535>
<0-2147483647>|<A.B.C.D> <0-65535>}
```

Example

Create an extended community list based on the route target attribute:

```
Switch(config)# ip extcommunity-list 1 memberid 234 rt 192.0.2.1 5 soo 32
45
```

Variable Definitions

Use the data in the following table to use the **ip extcommunity-list** command.

Variable	Value
<1-1024>	Specifies an integer value from 1–1024 that represents the community list ID you want to create or modify.
memberId <0-65535>	Specifies an integer value from 0–65535 that represents the member ID in the community list.
rt <0-65536> <0-2147483647> rt <A.B.C.D> <0-65535>	Specifies the route target in the format {AS number:assigned number} (that is, {0–65535}:{0–2147483647}) or {ipaddress:assigned number} (that is, {a.b.c.d}:{0–65535}).
soo <0-65535> <0-2147483647> soo <A.B.C.D> <0-65535>	Specifies the site of origin in the format {AS number:assigned number} (that is, {0–65535}:{0–2147483647}) or {ipaddress:assigned number} (that is, {a.b.c.d}:{0–65535}).

Chapter 5: BGP verification using CLI

Use **show** commands to verify Border Gateway Protocol (BGP) configuration and to monitor or troubleshoot BGP operation.

Viewing BGP aggregate information

Display information about current aggregate addresses.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about current aggregates:

```
show ip bgp aggregates [<prefix/len>] [vrf WORD <1-16>] [vrfids  
WORD<0-255>]
```

Variable definitions

Use the data in the following table to use the **show ip bgp aggregates** command.

Variable	Value
<prefix/len>	Specifies the IP address and the mask length.
vrf WORD<1-16>	Specifies a VRF instance by name.
vrfids WORD<0-255>	Specifies a range of VRFs by ID number.

Viewing IPv6 BGP+ aggregate information

Display information about current IPv6 aggregate addresses.

About this task

Use BGP 4 byte AS numbers to ensure the continuity of loop-free inter-domain routing information between ASs and to control the flow of BGP updates as 2 byte AS numbers will deplete soon.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about current IPv6 aggregates:

```
show bgp ipv6 aggregates [<WORD 1-256>] [vrf <WORD 1-16>] [vrfids
<0-255>]
```

Variable definitions

Use the data in the following table to use the **show bgp ipv6 aggregates** command.

Variable	Value
WORD <1-256>	Specifies the IPv6 prefix and the prefix length (the length can be 0 to 128).
vrf WORD <1-16>	Specifies a VRF instance by name (the string length ranges from 1-16 characters).
vrfids <0-255>	Specifies a range of VRFs by ID number (the ID ranges from 0-255).

Viewing CIDR routes

Display information about classless interdomain routing (CIDR) routes.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about CIDR routes:

```
show ip bgp cidr-only [<prefix/len>] [vrf WORD<1-16>] [vrfids
WORD<0-512>]
```

Variable definitions

Use the data in the following table to use the **show ip bgp cidr-only** command.

Variable	Value
<prefix/len>	Specifies an exact match of the prefix. This variable is an IP address and an integer value from 0-32 in the format a.b.c.d/xx.
vrf WORD<1-16>	Specifies a VRF instance by name (the string length ranges from 1-16 characters).

Table continues...

Variable	Value
vrfids WORD<0–512>	Specifies a range of VRFs by ID number (the ID ranges from 0–512).

Job Aid

Use the data in the following table to understand the `show ip bgp cidr-only` command output.

Table 10: show ip bgp cidr-only field descriptions

Field	Description
NETWORK/MASK	Specifies the network IP address and exact mask length (must be an integer value from 0–32).
PEER REM ADDR	Specifies the IP address of the remote peer.
NEXTHOP ADDRESS	Specifies the IP address of the next hop.
ORG	Specifies the source of a route: - IGP — the route is interior to the originating AS that inserts this route into the BGP table (0 = IGP). - EGP — the route is learned through an Exterior Gateway Protocol (EGP) before it is inserted into the BGP table (1 = BGP). - Incomplete — the origin of the route is unknown or learned by some other means. For example, the router learns these routes through RIP, OSPF, or static routes (2 = Incomplete).
LOC PREF	Specifies the local preference.

Viewing BGP configuration

View information about the BGP configuration.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about the current BGP configuration:

```
show ip bgp conf [vrf WORD<1-16>] [vrfids WORD<0-512>]
```

Example

```
=====
BGP Configuration - VRF vrf1
=====

BGP version - 4
```

```

        local-as - 22610
        Identifier - 27.82.217.1
        BGP on/off - ON
        as-4-byte - disable
        as-dot - disable
        aggregation - enable
        always-cmp-med - disable
        auto-peer-restart - enable
        auto-summary - enable
        comp-bestpath-med-confed - disable
        default-local-preference - 100
        default-metric - -1
        deterministic-med - disable
        flap-dampening - disable
        debug-screen - Off
        global-debug - none
        ibgp-report-import-rt - enable
        ignore-illegal-rtrid - enable
        max-equalcost-routes - 1
        no-med-path-is-worst - enable
        route-refresh - disable
        orig-def-route - disable
        orig-v6-def-route - disable
        quick-start - disable
        synchronization - enable

--More-- (q = quit)

```

Variable definitions

Use the data in the following table to use the **show ip bgp conf** command.

Variable	Value
vrf WORD<1–16>	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids WORD<0–512>	Specifies a range of VRFs by ID number (the ID ranges from 0–512).

Viewing BGP confederation

Display information about BGP confederations.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about current BGP confederations:

```
show ip bgp confederation
```

Example

```
Switch(config)#show ip bgp confederation
confederation identifier 0
confederation peer as
```

Viewing flap-dampened routes

Display information about flap-dampened routes to determine unreliable routes.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about flap-dampened routes:

```
show ip bgp dampened-paths {A.B.C.D} [<prefix/len>] [longer-
prefixes] [vrf WORD<1-16>] [vrfids WORD<0-512>]
```

Variable definitions

Use the data in the following table to use the **show ip bgp dampened-paths** command.

Variable	Value
{A.B.C.D}	Specifies the source IP address in the format a.b.c.d.
longer-prefixes	Shows long prefixes. The longer-prefixes indicate the mask length from a specified prefix to 32 (for example, show from prefix a.b.c.d/len to a.b.c./32).
<prefix/len>	Shows paths with this prefix. The prefix is the IP address and exact mask length (must be an integer value from 0–32).
vrf WORD<1–16>	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids WORD<0–512>	Specifies a range of VRFs by ID number (the ID ranges from 0–512).

Job Aid

Use the data in the following table to understand the **show ip bgp dampened-paths** command output.

Table 11: show ip bgp dampened-paths field descriptions

Field	Description
NETWORK/MASK	Specifies the network IP address and exact mask length (must be an integer value from 0–32).
PEER REM ADDR	Specifies the IP address of the remote peer.
NEXTHOP ADDRESS	Specifies the IP address of the next hop.
ORG	Specifies the source of a route: - IGP — the route is interior to the originating AS that inserts this route into the BGP table (0 = IGP). - EGP — the route is learned through an Exterior Gateway Protocol (EGP) before it is inserted into the BGP table (1 = BGP). - Incomplete — the origin of the route is unknown or learned by some other means. For example, the router learns these routes through RIP, OSPF, or static routes (2 = Incomplete).
LOC PREF	Specifies the local preference.

Viewing global flap-dampening configurations

Display global information about flap-dampening.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display global information about flap-dampening:

```
show ip bgp flap-damp-config [prefix/len] [vrf WORD<1-16>] [vrfids
WORD<0-512>]
```

Example

```
Switch(config)# show ip bgp flap-damp-config vrf vrf1
```

```
=====
                        BGP Flap Dampening - VRF vrf1
=====
                        Status - enable
                        PolicyName - N/A
                        CutoffThreshold - 1536
                        ReuseThreshold - 512
                        Decay - 2
                        MaxHoldDown - 180
```

Variable definitions

Use the data in the following table to use the **show ip bgp flap-damp-config** command.

Variable	Value
<prefix/len>	Shows paths with this prefix. The prefix is the IP address and exact mask length (must be an integer value from 0–32).
vrf WORD<1–16>	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids WORD<0–512>	Specifies a range of VRFs by ID number (the ID ranges from 0–512).

Job aid

Use the data in the following table to understand the **show ip bgp flap-damp-config** command output.

Table 12: show ip bgp flap-damp-config field descriptions

Field	Description
Status	Indicates the global state of the route flap dampening feature. Valid values are enable or disable.
PolicyName	This field does not apply to the switch.
CutoffThreshold	Indicates the penalty level that causes route suppression.
ReuseThreshold	Specifies the system-configured time for route reuse.
Decay	Indicates the decay rate based on the decay algorithm.
MaxHoldDown	Indicates the maximum length of time, in seconds, to suppress the route.

Viewing imported routes

Display information about BGP imported routes.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about BGP imported routes:

```
show ip bgp imported-routes [<prefix/len>] [longer-prefixes] [vrf
WORD<1–16>] [vrfids WORD<0–512>]
```

Variable definitions

Use the data in the following table to use the **show ip bgp imported-routes** command.

Variable	Value
longer-prefixes	Shows long prefixes. The longer-prefixes indicate the mask length from a specified prefix to 32 (for example, show from prefix a.b.c.d/len to a.b.c./32).
<prefix/len>	Shows paths with this prefix. The prefix is the IP address and exact mask length (must be an integer value from 0–32).
vrf WORD<1–16>	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids WORD<0–512>	Specifies a range of VRFs by ID number (the ID ranges from 0–512).

Job Aid

Use the data in the following table to understand the **show ip bgp imported-routes** command output.

Table 13: show ip bgp imported-routes field descriptions

Field	Description
ROUTE	Specifies the IP address of the route.
METRIC	Specifies the route metric.
COMMUNITY	Specifies the BGP community.
LOCALPREF	Specifies the local preference.
NEXTHOP	Specifies the IP address of the next hop.
ORIGIN	Specifies the source of a route: • IGP — the route is interior to the originating AS that inserts this route into the BGP table (0 = IGP). • EGP — the route is learned through an Exterior Gateway Protocol (EGP) before it is inserted into the BGP table (1 = BGP). • Incomplete — the origin of the route is unknown or learned by some other means. For example, the router learns these routes through RIP, OSPF, or static routes (2 = Incomplete).

Viewing BGPv6 imported routes

Display information about BGPv6 imported routes.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about BGPv6 imported routes:

```
show bgp ipv6 imported-routes [<prefix/len>] [longer-prefixes] [vrf
WORD<1-16>] [vrfids WORD<0-255>]
```

Variable definitions

Use the data in the following table to use the **show bgp ipv6 imported-routes** command.

Variable	Value
longer-prefixes	Shows long prefixes. The longer-prefixes indicate the mask length from a specified prefix to 32 (for example, show from prefix a.b.c.d/len to a.b.c./32).
<prefix/len>	Shows paths with this prefix. The prefix is the IP address and exact mask length (must be an integer value from 0–32).
vrf WORD<1–16>	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids WORD<0–255>	Specifies a range of VRFs by ID number (the ID ranges from 0–255).

Viewing BGP neighbors information

Display information about BGP neighbors.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about BGP neighbors:

```
show ip bgp neighbors [{A.B.C.D}] [vrf WORD<1-16>] [vrfids WORD<0-
512>]
```

3. Display information about BGP peer advertised routes:

```
show ip bgp neighbors {A.B.C.D} advertised-routes [<prefix/len>]
[longer-prefixes] [vrf WORD<1-16>] [vrfids WORD<0-512>]
```

4. Display information about BGP peer routes:

```
show ip bgp neighbors {A.B.C.D} routes [<prefix/len>] [community
<enable|disable>] [longer-prefixes] [vrf WORD<1-16>] [vrfids WORD<0-
512>]
```

5. Display statistics for BGP peers:

```
show ip bgp neighbors {A.B.C.D} stats [vrf WORD<1-16>] [vrfids
WORD<0-512>]
```

Example

```
Switch:#show ip bgp neighbors vrf vrf1
```

```
=====
BGP Neighbor Info - VRF vrf1
=====
```

```
BGP neighbor is 200.200.200.63 remote AS 63, Internal Peer, MP-BGP-capable, BGP state
[Established] UP Time 0 day(s), 07:27:24 remote router ID 63.1.1.1
```

```

vrf instance - 0
admin-state - BGP ON
connect-retry-interval - 120
ebgp-multihop - disable
hold-time - 30
keepalive-time - 10
hold-time-configured - 180
keepalive-time-configured - 60
max-prefix - 12000
nexthop-self - disable
originate-def-route - disable
MD5-authentication - disable
neighbor-debug - all
remove-private-as - disable
route-advertisement-interval - 5
route-reflector-client - disable
send-community - disable
soft-reconfiguration-in - disable
updt-source-interface - 0.0.0.0
weight - 100
Route Policy In -
Route Policy Out -
```

```

address-family vpnv4 - disable
route-refresh - disable
```

```
Total bgp neighbors - 1
```

Variable definitions

Use the data in the following table to use the **show ip bgp neighbors** command.

Variable	Value
{A.B.C.D}	Specifies the IP address.

Table continues...

Variable	Value
community <enable disable>	Enables or disables the display of community attributes.
longer-prefixes	Shows long prefixes. The longer-prefixes indicate the mask length from a specified prefix to 32 (for example, show from prefix a.b.c.d/len to a.b.c./32).
prefix/len	Shows paths with this prefix. The prefix is the IP address and exact mask length (must be an integer value from 0–32).
vrf WORD<1–16>	Specifies a VRF instance by name.
vrfids WORD<0–512>	Specifies a range of VRFs by ID number.

Viewing BGPv6 neighbors information

View information about BGPv6 neighbors.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. View information about BGPv6 neighbors:

```
show bgp ipv6 neighbors WORD<1-256> [vrf <WORD 1-16>] [vrfids
<0-255>]
```

3. View information about BGPv6 peer advertised routes:

```
show bgp ipv6 neighbors WORD<1-256> advertised-routes [WORD<1-256>]
[longer-prefixes] [vrf <WORD 1-16>] [vrfids <0-255>]
```

4. View information about BGPv6 peer routes:

```
show bgp ipv6 neighbors WORD<1-256> routes [WORD<1-256>] [community
<enable|disable>] [vrf <WORD 1-16>] [vrfids <0-255>]
```

Example

The following examples shows the summary output for **show bgp ipv6 neighbors** command, and the *advertised-routes* and *routes* variable options.

```
Switch:1>show bgp ipv6 neighbors vrf vrf1
```

```
=====
                        BGPv6 Neighbor Info - VRF vrf1
=====
BGPv6 neighbor is 2015:cdba:0:0:0:0:3257:9652 remote AS 200, External Peer,
BGP state [Established]
remote router ID 0.0.0.6

                        vrf instance - 0
                        admin-state - BGP ON
connect-retry-interval - 120
ebgp-multihop - disable
                        hold-time - 180
keepalive-time - 60
```

```

        hold-time-configured - 180
        keepalive-time-configured - 60
        ipv6-max-prefix - 8000
        nexthop-self - disable
        originate-defv6-route - disable
        neighbor-debug - all
        remove-private-as - disable

        route-advertisement-interval - 5
        route-reflector-client - disable
        send-community - disable
    soft-reconfiguration-in - enable
    updt-source-interface - 0:0:0:0:0:0:0
        weight - 100
    IPv6Route Policy In -
    IPv6Route Policy Out -
    address-family ipv6 - enable
    route-refresh - enable

Total bgpv6 neighbors: 1

Switch:1>show bgp ipv6 neighbors 2015:cdba:0:0:0:0:3257:9655 advertised-routes vrf vrf1
-----
The total number of routes advertised to the neighbor is 2
=====
                        BGPv6 Neighbor Advertised Routes - VRF vrf1
=====
NETWORK/MASK          NEXTHOP ADDRESS          LOC PREF ORG   STATUS
-----
2001:cdba:0:0:0:0:0/64  2001:cdba:0:0:0:0:3257:9651  100      INC   Best
2007:cdba:0:0:0:0:0/64  2001:cdba:0:0:0:0:3257:9651  100      INC   Used
-----

Switch:1>show bgp ipv6 neighbors 2015:cdba:0:0:0:0:3257:9655 routes vrf vrf1
-----
The total number of accepted routes from the neighbor is 2
=====
                        BGPv6 Neighbor Routes - VRF vrf1
=====
NETWORK/MASK          PEER-REM-ADDR          NEXTHOP-ADDRESS          ORG LOC-PREF STATUS
-----
1100:0:0:0:0:0:0/64    2015:cdba:0:0:0:0:3257:9655  2015:cdba:0:0:0:0:3257:9655  INC 100 Used AS_PATH:
(150)
2015:cdba:0:0:0:0:0:0/64  2015:cdba:0:0:0:0:3257:9655  2015:cdba:0:0:0:0:3257:9655  INC 100 Best AS_PATH:
(150)
-----

```

Variable definitions

Use the data in the following table to use the **show bgp ipv6 neighbors** command.

Variable	Value
WORD<1–256>	Specifies the IPv6 address.
advertised-routes	Specifies an IPv6 neighbors advertised routes.
routes	Specifies an IPv6 neighbors routes.
WORD<1–256>	Specifies an IPv6 address/length.

Table continues...

Variable	Value
longer-prefixes	Shows long prefixes. The longer-prefixes indicate the mask length from any specified prefix to 128. For example, show from prefix :X::X:X/len to X:X::X:X/ 128.
community <enable disable>	Enables or disables the display of community attributes.
vrf	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids	Specifies a range of VRFs by ID number (the ID ranges from 0–255).

Viewing BGP network configurations

Display information about BGP network configurations.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about BGP network configurations:

```
show ip bgp networks [<prefix/len>] [vrf WORD<1-16>] [vrfids WORD<0-512>]
```

Variable definitions

Use the data in the following table to use the **show ip bgp networks** command.

Variable	Value
<prefix/len>	Shows paths with this prefix. The prefix is the IP address and exact mask length (must be an integer value from 0–32).
vrf WORD<1–16>	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids WORD<0–512>	Specifies a range of VRFs by ID number (the ID ranges from 0–512).

Viewing IPv6 BGP+ network configurations

Display information about BGP+ network configurations.

Procedure

1. Log on to the switch to enter User EXEC mode.

2. Display information about BGP+ network configurations:

```
show bgp ipv6 networks <WORD 1-256> [vrf <WORD 1-16>] [vrfids
<0-255>]
```

Variable definitions

Use the data in the following table to use the **show bgp ipv6 networks** command.

Variable	Value
<WORD 1-256>	Specifies the IPv6 prefix and the prefix length (must be an integer value between 0 and 128).
vrf	Specifies a VRF instance by name (the string length ranges from 1-16 characters).
vrfids	Specifies a range of VRFs by ID number (the ID ranges from 0-255).

Viewing BGP peer group information

Display information about BGP peer groups.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about BGP peer groups:

```
show ip bgp peer-group [WORD<0-1536>] [vrf WORD<1-16>] [vrfids
WORD<0-512>]
```

Variable definitions

Use the data in the following table to use the **show ip bgp peer-group** command.

Variable	Value
vrf WORD<1-16>	Specifies a VRF instance by name (the string length ranges from 1-16 characters).
vrfids WORD<0-512>	Specifies a range of VRFs by ID number (the ID ranges from 0-512).
WORD<0-1536>	Specifies the name of the peer group (the string length ranges from 0-1536 characters).

Viewing BGP redistributed routes

Display information about BGP redistributed routes.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about BGP redistributed routes:

```
show ip bgp redistributed-routes [<prefix/len>] [vrf WORD<1-16>]  
[vrfids WORD<0-512>]
```

Variable definitions

Use the data in the following table to use the **show ip bgp redistributed-routes** command.

Variable	Value
<prefix/len>	Shows paths with this prefix. The prefix is the IP address and exact mask length (must be an integer value from 0–32).
vrf WORD<1–16>	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids WORD<0–255>	Specifies a range of VRFs by ID number (the ID ranges from 0–255).

Job aid

Use the data in the following table to understand the **show ip bgp redistributed-routes** command output.

Table 14: show ip bgp redistributed-routes field descriptions

Field	Description
SRC-VRF	Indicates the redistribution source VRF instance.
SRC	Indicates the redistribution source: RIP, Local, Static, or OSPF.
MET	Indicates the metric value.
MET-TYPE	Indicates the redistribution metric type.
ENABLE	Indicates whether the redistribution policy is enabled (T) true or disabled (F) false.
RPOLICY	The route policy currently assigned to the redistribution.

Viewing BGPv6 redistributed routes

Display information about BGPv6 redistributed routes.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about BGPv6 redistributed routes:

```
show bgp ipv6 redistributed-routes [vrf <WORD 1-16>] [vrfids <0-255>]
```

Variable Definitions

Use the data in the following table to use the **show bgp ipv6 redistributed-routes** command.

Variable	Value
vrf	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids	Specifies a range of VRFs by ID number (the ID ranges from 0–255).

Viewing a summary of BGP configurations

Display summarized information about BGP.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display summarized information about BGP:

```
show ip bgp summary [vrf WORD<1-16>] [vrfids WORD<0-512>]
```

Example

The following example shows partial output for the **show ip bgp summary** command.

```
Switch:1>show ip bgp summary vrf vrf1
```

```
=====
                        BGP Summary - VRF vrf1
=====

                        BGP version - 4
                        local-as - 22610
                        Identifier - 27.82.217.1
                        Decision state - Idle
The total number of routes is 0
```

```

BGP NEIGHBOR INFO :
  NEIGHBOR      RMTAS      STATE      HLDTM      KPALV      HLDCFG      KPCFG      WGHT      CONRTY      ADVINT      UPTIME
-----
192.0.2.1      22620      Active      0           0           180         60        100       120        5         0 day(s), 07:25:09
Total bgp neighbors: 1

BGP CONFEDERATION INFO :
confederation identifier 0
confederation peer as

--More-- (q = quit)

```

Variable definitions

Use the data in the following table to use the **show ip bgp summary** command.

Variable	Value
vrf WORD<1–16>	Specifies a VRF instance by name.
vrfids WORD<0–512>	Specifies a range of VRFs by ID number.

Job aid

Use the data in the following table to understand the **show ip bgp summary** command output.

Table 15: Variable definitions

Field	Description
BGP version	Specifies the version of BGP that runs on the router.
local-as	Specifies the local autonomous system number.
Identifier	Specifies the BGP identifier.
Decision state	Specifies the BGP process state.
NEIGHBOR	Specifies the IP address of the remote peer.
RMTAS	Specifies the AS number of the remote peer.
STATE	Specifies the peer operating state: Idle, Accept, Connect, Open, Open-sent, and Established.
HLDTM	Specifies the negotiated hold time timer.
KPALV	Specifies the keep alive timer.
HLDCFG	Specifies the configured hold time timer.
KPCFG	Specifies the configured keep alive timer.
WGHT	Specifies the weight value assigned to the peer.
CONRTY	Specifies the retry timer.
ADVINT	Specifies the advertisement interval.

Table continues...

Field	Description
UPTime	Specifies how long (in seconds) this peer has been in the established state, or how long since this peer was last in the established state. It is set to zero when a new peer is configured or when the router is booted. If the peer never reaches the established state, the value remains zero.

Viewing a summary of BGPv6 configurations

View a summary of BGP peering over IPv6 transport.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. View BGPv6 summary:

```
show bgp ipv6 summary [vrf <WORD 1-16>] [vrfids <0-255>]
```

Example

The following example shows partial output for the **show bgp ipv6 summary** command.

```
Switch:1>show bgp ipv6 summary vrf vrf1

=====
BGP ipv6 Summary - VRF vrf1
=====

BGP version - 4
  local-as - 200
  Identifier - 0.0.0.6
  Decision state - Idle
The total number of routes is 1

BGPv6 NEIGHBOR INFO :

NEIGHBOR          RMTAS  STATE      HLDTM  KPALV  HLD CFG  KPCFG  WGHT  CONRTY  ADVINT
-----
2001:DB8:0:0:0:0:ffff      50    Established 180    60    180    60    100  120    5

Total bgpv6 neighbors: 1

BGP CONFEDERATION INFO :
confederation identifier 0
confederation peer as

BGPv6 NETWORK INFO :

=====
BGPv6 Networks - VRF vrf1
=====
)
```

Variable Definitions

Use the data in the following table to use the **show bgp ipv6 summary** command.

Variable	Value
vrf	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids	Specifies a range of VRFs by ID number (the ID ranges from 0–255).

Viewing BGP routes

Display information about BGP routes.

*** Note:**

BGP stores route information on the AVL tree and this command retrieves that information. Information in the AVL tree is not sorted. The information returned by this command will not be displayed in any particular order.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Display information about BGP routes:

```
show ip bgp route [<prefix/len>] [community <enable|disable>] [ip
{A.B.C.D}] [longer-prefixes] [vrf WORD<1-16>] [vrfids WORD<0-512>]
```

Variable definitions

Use the data in the following table to use the **show ip bgp route** command.

Variable	Value
community <enable disable>	Enables or disables the display of community attributes.
ip {A.B.C.D}	Specifies an IP address.
longer-prefixes	Shows long prefixes. Longer-prefixes indicates the mask length from a specified prefix to 32 (for example, show from prefix a.b.c.d/len to a.b.c./32).
<prefix/len>	Shows paths with this prefix. The prefix is the IP address and exact mask length (must be an integer value from 0–32).
vrf WORD<1–16>	Specifies a VRF instance by name.
vrfids WORD<0–512>	Specifies a range of VRFs by ID number.

Job Aid

Use the data in the following table to understand the **show ip bgp route** command output.

Table 16: show ip bgp route

Field	Description
NETWORK/MASK	Specifies the path prefix address.
PEER REM ADDR	Specifies the remote peer address.
NEXTHOP ADDRESS	Specifies the BGP next hop address.
ORG	Specifies the source of a route: • IGP — the route is interior to the originating AS that inserts this route into the BGP table (0 = IGP). • EGP — the route is learned through an Exterior Gateway Protocol (EGP) before it is inserted into the BGP table (1 = BGP). • Incomplete — the origin of the route is unknown or learned by some other means. For example, the router learns these routes through RIP, OSPF, or static routes (2 = Incomplete).
LOCAL PREF	Specifies the local preference.

Viewing BGPv6 routes

Display information about BGPv6 routes.

Procedure

1. Log on to the switch to enter User EXEC mode.
2. Enter Privileged EXEC mode:

```
enable
```

3. Display information about BGP routes:

```
show bgp ipv6 route [<WORD 1-256> [longer-prefixes]] [community
<enable|disable>] [ipv6 <WORD 1-256>] [vrf <WORD 1-16>] [vrfs
<0-255>]
```

Variable definitions

Use the data in the following table to use the **show bgp ipv6 route** command.

Variable	Value
[<WORD 1-256>]	Specifies the IPv6 prefix and the prefix length (must be an integer value between 0 and 128).
community <enable disable>	Enables or disables the display of community attributes.
ipv6 <WORD 1-256>]	Specifies an IPv6 address.
longer-prefixes	Shows long prefixes. The longer-prefixes indicate the mask length from any specified prefix to 128 (for example, show from prefix X:X::X:X/len to X:X::X:X/ 128).
vrf	Specifies a VRF instance by name (the string length ranges from 1–16 characters).
vrfids	Specifies a range of VRFs by ID number (the ID ranges from 0–255).

Chapter 6: BGP configuration using EDM

Configure Border Gateway Protocol (BGP) to create an inter-domain routing system that guarantees loop-free routing information between autonomous systems.

For information about how to configure route policies, see [Configuring IPv4 Routing for VOSS](#).

Configuring BGP globally

Enable BGP so that BGP runs on the router. Configure general BGP parameters to define how BGP operates on the system.

Before you begin

- Change the VRF instance as required to configure BGP on a specific VRF instance. The VRF must have an RP trigger of BGP. Not all parameters are configurable on non-default VRFs.

About this task

If you must configure the BGP router ID, use CLI. You cannot configure the BGP router ID using EDM.

Procedure

1. In the navigation pane, expand the following folders: **Configuration > IP**.
2. Click **BGP**.
3. Click the **Generals** tab.
4. In AdminStatus, select **enable**.
5. Configure the local autonomous system (AS) ID.
6. In the **Aggregate** area, enable or disable route aggregation as required.
7. Configure the BGP options as required.
8. In the **DebugMask** area, select the check box for the type of information to show for BGP debugging purposes.
9. Configure BGP confederations as required.
10. Configure BGP route reflectors as required.
11. Click **Apply**.

Generals Field Descriptions

Use the data in the following table to use the **Generals** tab.

Name	Description
bgpVersion	Specifies the version of BGP that operates on the router.
bgpIdentifier	Specifies the BGP router ID number.
AdminStatus	Enables or disables BGP on the router. The default is disable. You cannot enable AdminStatus until you change the LocalAS value to a nonzero value.
4ByteAs	Enables or disables 4-byte AS numbers. The default is disable.
LocalAs	Configures the local AS number in the range of 0–65535. You cannot change the LocalAS if AdminStatus is enable.
AsDot	Enables or disable the AS dot notation format for the 4-byte AS number. The default is disable. The AS dot notation is easier to read and remember than the AS plain notation, but it can be difficult to convert from AS plain to AS dot. The IETF prefers the AS plain notation.
Aggregate	Enables or disables aggregation. The default is enable.
DefaultMetric	Configures the metric sent to BGP neighbors. The default metric determines the cost of a route a neighbor uses. Use this parameter in conjunction with the redistribute parameters so that BGP uses the same metric for all redistributed routes. The default is -1. The range is -1–2147483647.
DefaultLocalPreference	Specifies the default local preference. The local preference indicates the preference that AS border routers assign to a chosen route when they advertise it to iBGP peers. The default is 100. The range is 0–2147483647.
AlwaysCompareMed	Enables or disables the comparison of the multi-exit discriminator (MED) parameter for paths from neighbors in different autonomous systems. The system prefers a path with a lower MED over a path with a higher MED. The default is disable.
DeterministicMed	Enables or disables deterministic MED. Deterministic MED compares the MEDs after routes advertised by different peers in the same AS are chosen. The default is disable.
AutoPeerRestart	Enables or disables the process that automatically restarts a connection to a BGP neighbor. The default is enable.
AutoSummary	Enables or disables automatic summarization. If you enable this variable, BGP summarizes networks based on class limits (for example, Class A, B, or C networks). The default is enable.

Table continues...

Name	Description
NoMedPathsWorst	Enables or disables NoMedPathsWorst. If you enable this variable, BGP treats an update without a MED attribute as the worst path. The default is enabled.
BestPathMedConfed	Enables or disables the comparison of MED attributes within a confederation. The default is disable.
DebugMask	Displays the specified debug information for BGP global configurations. The default value is none. Other options are • none disables all debug messages. • event enables the display of debug event messages. • state enables display of debug state transition messages. • update enables display of debug messages related to updates transmission and reception. • error enables the display of debug error messages. • trace enables the display of debug trace messages. • init enables the display of debug initialization messages. • all enables all debug messages. • packet enables the display of debug packet messages. • warning enables the display of debug warning messages. • filter enables the display of debug messages related to filtering.
IgnoreIllegalRouterId	Enables BGP to overlook an illegal router ID. For example, this variable enables the acceptance of a connection from a peer that sends an open message using a router ID of 0. The default is enable.
Synchronization	Enables or disables the router to accept routes from BGP peers without waiting for an update from the IGP. The default is enable.
MaxEqualcostRoutes	Configures the maximum number of equal-cost-paths that are available to a BGP router by limiting the number of equal-cost-paths the routing table can store. The default value is 1; the range is 1–8.
IbgpReportImportRoute	Configures BGP to report imported routes to an interior BGP (iBGP) peer. This variable also enables or disables reporting of non-BGP imported routes to other iBGP neighbors. The default is enable.
FlapDampEnable	Enables or disables route suppression for routes that go up and down (flap). The default is disable.

Table continues...

Name	Description
QuickStart	Enables or disables the Quick Start feature, which forces the BGP speaker to begin establishing peers immediately, instead of waiting for the auto-restart timer to expire. The default is disable.
TrapEnable	Enables or disables the BGP traps. The default is disable.
ConfederationASIdentifier	Specifies a BGP confederation identifier in the range of 0–65535. * Note: This parameter applies only to VRF 0.
ConfederationPeers	Lists adjoining autonomous systems that are part of the confederation in the format (5500,65535,0,10,...,...).. This value can use 0–255 characters. * Note: This parameter applies only to VRF 0.
RouteReflectionEnable	Enables or disables the reflection of routes from iBGP neighbors. The default is enable. * Note: This parameter applies only to VRF 0.
RouteReflectorClusterId	Configures a reflector cluster ID IP address. This variable applies only if you enable RouteReflectionEnable, and if multiple route reflectors are in a cluster. * Note: This parameter applies only to VRF 0.
ReflectorClientToClientReflection	Enables or disables route reflection between two route reflector clients. This variable applies only if RouteReflectionEnable is enable. The default is enable. * Note: This parameter applies only to VRF 0.
RouteRefresh	Enables or disables route refresh. If enabled, a route refresh request received by a BGP speaker causes the speaker to resend all route updates it contains in its database that are eligible for the peer that issues the request. This parameter only applies to VRF 0.

Configuring 4-byte AS numbers

Configure AS numbers using the 4-byte format and represent the numbers in octets.

Before you begin

- You cannot modify the global BGP configuration unless BGP is disabled.
- Make sure that you define AS numbers in policies the same way that you configure them for the router. The choices are asplain (regular expression) or asdot (dot notation). If you create policies using asplain and configure the switch with asdot, the match will not occur.

About this task

Use BGP 4–byte AS numbers to ensure the continuity of loop-free inter-domain routing information between autonomous systems and to control the flow of BGP updates as 2 byte AS numbers will deplete soon. AS Plain notation format is the default and the preferred form of representing 4–byte AS numbers over the AS dot notation format.

You have an option to configure AS dot notation format as well. With AS dot notation, analyzing and troubleshooting any issues encountered becomes difficult as it is incompatible with the regular expressions used by most of the network providers.

Procedure

1. In the navigation pane, open the following folders: **Configuration > IP**.
2. Click **BGP**.
3. Click the **Generals** tab.
4. To change the AS number format, select **disable** for **AdminStatus**.
5. Click **Apply**.
6. In **4-byteAs** , select **enable**.
7. In **AsDot**, select **enable**.
8. Enter the 4-byte AS number in octets in the **LocalAs** field.
9. In **AdminStatus**, select **enable**.
10. Click **Apply**.

4–byte AS field descriptions

Use the data in the following table to use the 4–byte AS related fields on the **Generals** tab.

Name	Description
LocalAs	Configures the local autonomous system (AS) number. You cannot change this field when AdminStatus is set to enable. This field sets a 2-byte local AS number in the range from 0 to 65535. To set a 4-byte local AS number, click enable in the 4ByteAs field and enter a number in the LocalAs field.

Table continues...

Name	Description
	Attention: The switch does not support this parameter with BGP +.
4byteAs	Enables or disables the switch from using 4 byte numbers for autonomous systems.
AsDot	Enables or disables representing AS numbers in octets. The default is disable so the switch uses the plain notation format. If you enable this field and the 4ByteAs field, enter the AS number in the LocalAs field. Attention: The switch does not support this parameter with BGP +.

Configuring aggregate routes

Configure aggregate routes so that the router advertises a single route (aggregate route) that represents all destinations. Aggregate routes also reduce the size of routing tables.

Before you begin

- Enable aggregate routes globally.
- You need the appropriate aggregate address and mask.
- If required, policies exist.
- Change the VRF instance as required to configure BGP on a specific VRF instance. The VRF must have an RP trigger of BGP. Not all parameters are configurable on non-default VRFs.

Procedure

1. In the navigation pane, expand the following folders: **Configuration > IP**.
2. Click **BGP**.
3. Click the **Aggregates** tab.
4. Click **Insert**.
5. Configure the aggregate **Address** and **PrefixLen**.
6. Select **AsSetGenerate** or **SummaryOnly** as required.
7. Configure policies for the aggregate route.
8. Click **Insert**.

Aggregates field descriptions

Use the data in the following table to use the **Aggregates** tab.

Name	Description
Address	Specifies the aggregate IP address.
PrefixLen	Specifies the aggregate PrefixLen.
AsSetGenerate	Enables or disables AS-set path information generation. The default is disable.
SummaryOnly	Enables or disables the summarization of routes in routing updates. Enable this parameter to create the aggregate route and suppress advertisements of more-specific routes to all neighbors. The default is disable.
SuppressPolicy	Specifies the route policy (by name) used for the suppressed route list. Enable this parameter to create the aggregate route and suppress advertisements of the specified routes.
AdvertisePolicy	Specifies the route policy (by name) used for route advertisements. The route policy selects the routes that create AS-set origin communities.
AttributePolicy	Specifies the route policy (by name) used to determine aggregate route attributes.

Configuring aggregate IPv6 routes

Configure IPv6 aggregate routes so that the router advertises a single route (aggregate route) that represents all destinations. Aggregate routes also reduce the size of routing tables.

To configure aggregate routes for IPv4, see [Configuring aggregate routes](#) on page 102.

Before you begin

- Aggregate routes are enabled.
- You have determined the appropriate aggregate prefix and length.
- If required, policies exist.

Procedure

1. In the navigation pane, expand the following folders: **Configuration > IPv6**.
2. Click **BGP+**.
3. Click the **Aggregates** tab.
4. Click **Insert**.
5. Specify the **aggregate Address** and **PrefixLen**
6. Configure **AsSetGenerate** and **SummaryOnly** as required.
7. Configure policies for the aggregate route.
8. Click **Insert**.

Aggregates field descriptions

Use the data in the following table to use the **Aggregates** tab.

Name	Description
Address	Specifies the aggregate address. The default is none.
PrefixLen	Specifies the length of the prefix (in bits).
AsSetGenerate	Enables or disables AS-set path information generation. The default is disable.
SummaryOnly	Enables or disables the summarization of routes in routing updates. Enable this parameter to create the aggregate route and suppress advertisements of more-specific routes to all neighbors. The default is disable.
SuppressPolicy	Specifies the route policy (by name) used for the suppressed route list. Enable this parameter to create the aggregate route and suppress advertisements of the specified routes.
AdvertisePolicy	Specifies the route policy (by name) used for route advertisements. The route policy selects the routes that create AS-set origin communities.
AttributePolicy	Specifies the route policy (by name) used to determine aggregate route attributes.

Configuring allowed networks

Configure network addresses to determine the network addresses that BGP advertises. The allowed addresses determine the BGP networks that originate from the switch.

Before you begin

- Change the VRF instance as required to configure BGP on a specific VRF instance. The VRF must have an RP trigger of BGP. Not all parameters are configurable on non-default VRFs.

Procedure

1. In the navigation pane, expand the following folders: **Configuration > IP**.
2. Click **BGP**.
3. Click the **Network** tab.
4. Click **Insert**.
5. Configure the network address, mask, and metric.

- Click **Insert**.

Network field descriptions

Use the data in the following table to use the **Network** tab.

Name	Description
NetworkAfAddr	Specifies the network prefix that BGP advertises.
NetworkAfPrefixLen	Specifies the prefix length of the network address.
NetworkAfMetric	Specifies the metric to use when the system sends an update for the routes in the network table. The metric configures the MED for the routes advertised to eBGP peers. The range is 0–65535.

Configuring allowed IPv6 networks

Configure IPv6 network addresses to determine the network addresses that BGP advertises. The allowed addresses determine the BGP networks that originate from the switch.

To configure allowed IPv4 networks, see [Configuring allowed networks](#) on page 104.

Procedure

- In the navigation pane, expand the following folders: **Configuration > IPv6**.
- Click **BGP+**.
- Click the **Network** tab.
- Click **Insert**.
- Configure the **network address**, **prefix length** and **metric**.
- Click **Insert**.

Network field descriptions

Use the data in the following table to use the **Network** tab.

Name	Description
NetworkAfAddr	Specifies the network prefix that BGP advertises. The default is none.
NetworkAfPrefixLen	Specifies the network prefix length. The default is none.

Table continues...

Name	Description
NetworkAfMetric	Specifies the metric used when an update is sent for the routes in the network table. The metric configures the MED for the routes advertised to EBGP peers. The range is 0 to 65535. The default is 0.

Configuring BGP peers

Configure BGP peers to connect two routers to each other for the purpose of exchanging routing information. BGP peers exchange complete routing information only after they establish the peer connection.

Before you begin

- Change the VRF instance as required to configure BGP on a specific VRF instance. The VRF must have an RP trigger of BGP. Not all parameters are configurable on non-default VRFs.

Procedure

1. In the navigation pane, expand the following folders: **Configuration > IP**.
2. Click **BGP**.
3. Click the **Peers** tab.
4. Click **Insert**.
5. Configure the peer as required.
6. Click **Insert**.
7. In the **Enable** column, double-click the value, and then select **enable**.
By default, new peer configuration parameters are disabled.
8. Click **Apply**.
9. To modify a peer configuration, double-click the value, and then select a new value.

Peers Field Descriptions

Use the data in the following table to use the **Peers** tab.

Name	Description
RemoteAddr	Specifies the remote IP address of the entered BGP peer.
GroupName	Specifies the peer group name to which the peer belongs (optional).
PeerState	Specifies the BGP peer connection state.

Table continues...

Name	Description
RemoteAs	Configures a remote AS number for the peer or peer-group in the range 0–65535.
Enable	Controls whether the peer connection is enabled or disabled. The default is disabled.
EbgpMultiHop	Enables or disables a connection to a BGP peer that is more than one hop away from the local router. The default value is disable.
RoutePolicyIn	Specifies the policy (by name) that applies to all network routes learned from this peer.
RoutePolicyOut	Specifies the policy (by name) that applies to all outgoing route updates.
RemovePrivateAs	Strips (when enabled) private AS numbers when the switch sends an update. The default is enable.
UpdateSourceInterface	Specifies the source IP address to use when the switch sends eBGP packets to this peer or peer group.
ConnectRetryInterval	Specifies the time interval, in seconds, for the connect retry timer. The suggested value for this timer is 120 seconds. The range is 1 to 65535.
HoldTimeConfigured	Specifies the time interval, in seconds, for the hold time for this BGP speaker with this peer. This value is in an open message sent to this peer by this BGP speaker. To determine the hold time with the peer, the switch compares this value with the HoldTime value in an open message received from the peer. The HoldTime must be at least three seconds. If the value is zero, the hold time does not establish with the peer. The suggested value for this timer is 180 seconds. The range is 0 to 65535.
KeepAliveConfigured	Specifies the time interval, in seconds, for the KeepAlive timer configured for this BGP speaker with this peer. KeepAliveConfigured determines the keep alive message frequency relative to HoldTimeConfigured; KeepAlive indicates the actual time interval for the keep alive messages. The recommended maximum value for this timer is one-third of HoldTimeConfigured. If KeepAliveConfigured is zero, no periodic keep alive messages are sent to the peer after the peers establish a BGP connection. Configure a value of 60 seconds. The range is 0 to 21845.
MD5Authentication	Enables and disables MD5 authentication.
AdvertisementInterval	Specifies the time interval, in seconds, that elapses between each transmission of an advertisement from a BGP neighbor. The default value is 30 seconds and the range is 5–120 seconds. The route advertisement interval feature is implemented using the time stamp that indicates when each route is advertised. The time stamp is marked to each route so that the route advertisement interval is compared to the time stamp and BGP is then able to make a decision about whether the route advertisement can be sent or it

Table continues...

Name	Description
	should be delayed when a better route is received. This feature does not work for a withdraw route because the route entry is already removed when the processing route advertisement is sent and the time stamp marked in the route entry cannot be obtained.
DefaultOriginate	When enabled, specifies that the current route originated from the BGP peer. This parameter enables or disables sending the default route information to the specified neighbor or peer. The default value is false.
DefaultOriginateIpv6	When enabled, specifies that the current IPv6 route originated from the BGP peer. This parameter enables or disables sending the default IPv6 route information to the specified neighbor or peer. The default value is false.
Weight	Specifies the peer or peer group weight, or the priority of updates the system can receive from this BGP peer. The default value is 100 and the range is 0–65535.
MaxPrefix	Configures a limit on the number of routes accepted from a neighbor. The default value is 12000 routes and the range is 0–2147483647. A value of 0 means no limit exists.
NextHopSelf	Specifies that the next-hop attribute in an iBGP update is the address of the local router or the router that generates the iBGP update. The default is disable.
RouteReflectorClient	Specifies that this peer is a route reflector client.  Note: This parameter only applies to VRF 0.
SoftReconfigurationIn	When enabled, the router relearns routes from the specified neighbor or group of neighbors without restarting the connection after the policy changes in the inbound direction. The default value is disable. Enabling SoftReconfigurationIn stores all BGP routes in local memory (even non-best routes).
DebugMask	Displays the specified debug information for the BGP peer. The default value is none. • None disables all debug messages. • Event enables the display of debug event messages. • State enables display of debug state transition messages. • Update enables display of debug messages related to updates transmission and reception. • Error enables the display of debug error messages. • Trace enables the display of debug trace messages. • Init enables the display of debug initialization messages.

Table continues...

Name	Description
	• All enables all debug messages. • Packet enables the display of debug packet messages. • Warning enables the display of debug warning messages. • Filter enables the display of debug messages related to filtering.
SendCommunity	Enables or disables sending the community attribute of the update message to the specified peer. The default value is disable.
IvpnLiteCap	Enable or disable IP VPN-lite capability on the BGP neighbor peer.
Ipv6Cap	Enable or disable the IPv6 capability on the BGP neighbor peer. The default value is disable.
RouteRefresh	Enables or disables route refresh. If enabled, a route refresh request received by a BGP speaker causes the speaker to resend all route updates it contains in its database that are eligible for the peer that issues the request. This parameter only applies to VRF 0.
AsOverride  Note: This field does not appear on all hardware platforms.	Specifies that the AS Override parameter can be enabled or disabled for the BGP peer. The default is disable.
AllowAsIn  Note: This field does not appear on all hardware platforms.	Specifies the number of AS-in allowed for the BGP peer. The range is 1–10.
Ipv6RoutePolicyIn	Specifies the policy (by name) that applies to all network IPv6 routes learned from this peer.
Ipv6RoutePolicyOut	Specifies the policy (by name) that applies to all outgoing IPv6 route updates.
Ipv6MaxPrefix	Configures a limit on the number of IPv6 routes accepted from a neighbor. A value of 0 means no limit exists.

Configuring BGPv6 peers

Configure BGPv6 peers to connect two routers to each other for the purpose of exchanging routing information. BGPv6 peers exchange complete routing information only after they establish the peer connection.

Procedure

1. In the navigation pane, expand the following folders: **Configuration > IPv6**.

2. Click **BGP+**.
3. Click the **Peers** tab.
4. Click **Insert**.
5. Configure the peer, as required.
6. Click **Insert**.
7. In the **Enable** column, double-click the value, and then select **enable**.
By default, new peer configuration parameters are disabled.
8. Click **Apply**.
9. To modify a peer configuration, double-click the value, and then select a new value.

Peers Field Descriptions

Use the data in the following table to use the Peers tab.

Name	Description
RemoteAddr	Specifies the remote IPv6 address of the entered BGP+ peer.
GroupName	Specifies the peer group name to which the peer belongs (optional).
PeerState	Specifies the BGPv6 peer connection state.
RemoteAs	Configures a remote AS number for the peer or peer-group in the range 0 to 65535.
Enable	Controls whether the peer connection is enabled or disabled. The default is disabled.
EbgpMultiHop	Enables or disables a connection to a BGPv6 peer that is more than one hop away from the local router. The default value is disable.
RoutePolicyIn	Specifies the policy (by name) that applies to all network routes learned from this peer.
RoutePolicyOut	Specifies the policy (by name) that applies to all outgoing route updates.
RemovePrivateAs	Strips (when enabled) private AS numbers when the switch sends an update. The default is enable.
UpdateSourceInterface	Specifies the source IP address to use when the switch sends eBGP packets to this peer or peer group.
ConnectRetryInterval	Specifies the time interval, in seconds, for the connect retry timer. The suggested value for this timer is 120 seconds. The range is 1 to 65535.
HoldTimeConfigured	Specifies the time interval, in seconds, for the hold time for this BGP speaker with this peer. This value is in an open message sent to this peer by this BGP speaker. To determine the hold time with the peer, the switch compares this value with the HoldTime value in an open

Table continues...

Name	Description
	message received from the peer. The HoldTime must be at least three seconds. If the value is zero, the hold time does not establish with the peer. The suggested value for this timer is 180 seconds. The range is 0 to 65535.
KeepAliveConfigured	Specifies the time interval, in seconds, for the KeepAlive timer configured for this BGP speaker with this peer. KeepAliveConfigured determines the keep alive message frequency relative to HoldTimeConfigured; KeepAlive indicates the actual time interval for the keep alive messages. The recommended maximum value for this timer is one-third of HoldTimeConfigured. If KeepAliveConfigured is zero, no periodic keep alive messages are sent to the peer after the peers establish a BGP connection. Configure a value of 60 seconds. The range is 0 to 21845.
MD5Authentication	Enables and disables MD5 authentication.
AdvertisementInterval	Specifies the time interval, in seconds, that elapses between each transmission of an advertisement from a BGPv6 neighbor. The default value is 30 seconds and the range is 5 to 120 seconds. The route advertisement interval feature is implemented using the time stamp that indicates when each route is advertised. The time stamp is marked to each route so that the route advertisement interval is compared to the time stamp and BGP is then able to make a decision about whether the route advertisement can be sent or it should be delayed when a better route is received. This feature does not work for a withdraw route because the route entry is already removed when the processing route advertisement is sent and the time stamp marked in the route entry cannot be obtained.
DefaultOriginatelpv6	When enabled, specifies that the current IPv6 route originated from the BGP peer. This parameter enables or disables sending the default IPv6 route information to the specified neighbor or peer. The default value is false.
Weight	Specifies the peer or peer group weight, or the priority of updates the system can receive from this BGP peer. The default value is 100 and the range is 0 to 65535.
MaxPrefix	Configures a limit on the number of routes accepted from a neighbor. The default value is 12000 routes and the range is 0 to 2147483647. A value of 0 means no limit exists.
NextHopSelf	Specifies that the next-hop attribute in an iBGP update is the address of the local router or the router that generates the iBGP update. The default is disable.
RouteReflectorClient	Specifies that this peer is a route reflector client.  Note: This parameter only applies to VRF 0.

Table continues...

Name	Description
SoftReconfigurationIn	When enabled, the router relearns routes from the specified neighbor or group of neighbors without restarting the connection after the policy changes in the inbound direction. The default value is disable. Enabling SoftReconfigurationIn stores all BGP routes in local memory (even non-best routes).
DebugMask	Displays the specified debug information for the BGP peer. The default value is none. • None disables all debug messages. • Event enables the display of debug event messages. • State enables display of debug state transition messages. • Update enables display of debug messages related to updates transmission and reception. • Error enables the display of debug error messages. • Trace enables the display of debug trace messages. • Init enables the display of debug initialization messages. • All enables all debug messages. • Packet enables the display of debug packet messages. • Warning enables the display of debug warning messages. • Filter enables the display of debug messages related to filtering.
SendCommunity	Enables or disables sending the community attribute of the update message to the specified peer. The default value is disable.
IvpnLiteCap	Enable or disable IP VPN-lite capability on the BGP neighbor peer.
Ipv6Cap	Enable or disable the IPv6 capability on the BGP neighbor peer. The default value is disable.
RouteRefresh	Enables or disables route refresh. If enabled, a route refresh request received by a BGP speaker causes the speaker to resend all route updates it contains in its database that are eligible for the peer that issues the request. This parameter only applies to VRF 0.
AsOverride  Note: This field does not appear on all hardware platforms.	Specifies that the AS Override parameter can be enabled or disabled for the BGP peer. The default is disable.
AllowAsIn  Note: This field does not appear on all hardware platforms.	Specifies the number of AS-in allowed for the BGP peer. The range is 1–10.

Table continues...

Name	Description
Ipv6RoutePolicyIn	Specifies the policy (by name) that applies to all network IPv6 routes learned from this peer.
Ipv6RoutePolicyOut	Specifies the policy (by name) that applies to all outgoing IPv6 route updates.
Ipv6MaxPrefix	Configures a limit on the number of IPv6 routes accepted from a neighbor. A value of 0 means no limit exists.

Configuring peer groups

Configure or edit peer groups to create update policies for neighbors in the same group.

Before you begin

- Change the VRF instance as required to configure BGP on a specific VRF instance. The VRF must have an RP trigger of BGP. Not all parameters are configurable on non-default VRFs.

Procedure

1. In the navigation pane, open the following folders: **Configuration > IP**.
2. Click **BGP**.
3. Click the **Peer Groups** tab.
You can modify an existing parameter by double-clicking the value.
4. Click **Insert**.
5. Configure the peer group as required.
6. Click **Insert**.

Peer Groups field descriptions

Use the data in the following table to use the **Peer Groups** tab.

Name	Description
Index	Specifies the index of this peer group.
GroupName	Specifies the peer group to which this neighbor belongs (optional).
Enable	Enables or disables the peer group.
RemoteAs	Configures a remote AS number for the peer-group in the range 0–65535.

Table continues...

Name	Description
DefaultOriginate	When enabled, the BGP speaker (the local router) sends the default route 0.0.0.0 to a group of neighbors for use as a default route. The default is disabled.
DefaultOriginateIpv6	When enabled, the BGP speaker (the local router) sends the default route to a group of neighbors for use as a default route. The default is disabled.
EbgpMultiHop	When enabled, the switch accepts and attempts BGP connections to external peers that reside on networks that do not directly connect. The default is disabled.
AdvertisementInterval	Specifies the time interval, in seconds, that elapses between BGP routing updates. The default value is 30 seconds.
KeepAlive	Specifies the time interval, in seconds, between sent BGP keep alive messages to remote peers. The default value is 60.
HoldTime	Configures the hold time for the group of peers in seconds. Use a value that is three times the value of the KeepAlive time. The default value is 180.
Weight	Assigns an absolute weight to a BGP network. The default value is 100.
MaxPrefix	Limits the number of routes accepted from this group of neighbors. A value of zero indicates no limit. The default value is 12,000 routes.
NextHopSelf	Specifies that the switch must set the NextHop attribute to the local router address before it sends updates to remote peers.
RoutePolicyIn	Specifies the route policy that applies to all networks learned from this group of peers.
RoutePolicyOut	Specifies the route policy that applies to all outgoing updates to this group of peers.
RouteReflectorClient	Specifies that this peer group is a route reflector client.  Note: This parameter only applies to VRF 0.
SoftReconfigurationIn	When enabled, the router relearns routes from the specified neighbor or group of neighbors without restarting the connection after the policy changes in the inbound direction. The default value is enable. Enabling SoftReconfigurationIn stores all BGP routes in local memory (even non-best routes).
MD5Authentication	Enables and disables MD5 authentication. The default is disable.
RemovePrivateAs	Strips (when enabled) private AS numbers when the switch sends an update. The default is enable.
SendCommunity	Enables or disables sending the community attribute of the update message to the specified peer group. The default value is disable.
AfUpdateSourceInterfaceType	Specifies the interface type.

Table continues...

Name	Description
AfUpdateSourceInterface	Specifies the IP address used for circuitless IP (CLIP) for this peer group.
Vpvnv4Address	Enables BGP address families for IPv4 (BGP) and L3 VPN (MP-BGP) support. Enable this parameter for VPN/VRF Lite routes.
IpvpnLiteCap	Specifies (when enabled) that IP VPN Lite capability can be enabled or disabled on the BGP neighbor peer. The default is disable.
RouteRefresh	Enables or disables route refresh. If enabled, a route refresh request received by a BGP speaker causes the speaker to resend all route updates it contains in its database that are eligible for the peer that issues the request. This parameter only applies to VRF 0.
AsOverride	Specifies that the AS Override parameter can be enabled or disabled for the BGP peer group. The default is disable.
AllowedAsIn	Specifies the number of AS-in allowed for the BGP peer group. The range is 1–10.
IPv6Cap	Enable or disable the IPv6 capability on the BGP neighbor peer. The default value is disable.
Ipv6RoutePolicyIn	Specifies the policy (by name) that applies to all network IPv6 routes learned from this peer.
Ipv6RoutePolicyOut	Specifies the policy (by name) that applies to all outgoing IPv6 route updates.
Ipv6MaxPrefix	Configures a limit on the number of IPv6 routes accepted from a neighbor. A value of 0 means no limit exists.

Viewing IPv6 Community Attributes

View IPv6 community attributes for specific routes to utilize the update message fields to communicate information between BGP speakers. The Path Attribute values allow you to specify the prefixes that the BGP session can exchanged, or which of the multiple paths of a specified prefix to use.

Procedure

1. In the navigation tree, open the following folders: **Configuration > IPv6**.
2. Click **BGP+**.
3. Click the **Bgp Route Summary** tab.
4. Select a route for which you want to view the route summary information.
5. Click the **Route Comm Attr** option on the menu.

The **BGP Path Attributes** tab opens with the BGP IPv6 community attribute information.

BGP Path Attributes field descriptions

Use the data in the following table to use the **Community List** tab.

Name	Description
Origin	Specifies the ultimate origin of the path information.
NextHopAddr	Specifies the address of the border router that is used to access the destination network. This address is the nexthop address received in the UPDATE packet associated with this prefix.
Med	This metric is used to discriminate between multiple exit points to an adjacent autonomous system. When the MED value is absent but has a calculated default value, this object will contain the calculated value.
LocalPref	Specifies the value used during route decision process in the BGP protocol. Applicable to BGP only.
AggregatorAS	Specifies the AS number of the last BGP4 speaker that performed route aggregation. If the AGGREGATOR path attribute is absent, this object will not be present in the conceptual row.
AggregatorAddr	Specifies the IP address of the last BGP4 speaker that performed route aggregation. If the AGGREGATOR path attribute is absent, this object will not be present in the conceptual row.
String	This is a string representing the autonomous system path to the network which was received from the peer which advertised it. The format of the string is implementation-dependent, and is designed for operator readability.  Note: SnmpAdminString is only capable of representing a maximum of 255 characters. This may lead to the string being truncated in the presence of a large AS Path.

Displaying dampened routes information

Display dampened path information to see which routes are suppressed.

Before you begin

- Change the VRF instance as required to view BGP information about a specific VRF instance. The VRF must have an RP trigger of BGP. Not all parameters are configurable on non-default VRFs.
- Enable dampened routes.

Procedure

1. In the navigation pane, open the following folders: **Configuration > IP**.
2. Click **BGP**.
3. Click the **Dampened Routes** tab.

Dampened Routes field descriptions

Use the data in the following table to use the **Dampened Routes** tab.

Name	Description
IpAddrPrefix	Specifies the IP address prefix in the NLRI field. This variable is an IP address that contains the prefix with a length specified by IpAddrPrefixLen. Bits beyond the length specified by IpAddrPrefixLen are set to zero.
IpAddrPrefixLen	Specifies the length, in bits, of the IP address prefix in the NLRI field.
Peer	Specifies the IP address of the peer from which the router learns the path information.
FlapPenalty	Specifies the penalty based on number of route flaps.
FlapCount	Specifies the number of times a route flapped (went down and up) since the last time the penalty was reset to zero.
RouteDampened	Indicates whether this route is suppressed or announced.
ReuseTime	Specifies the system-configured time for route reuse.

Configuring redistribution to BGP

Configure redistribute entries for BGP to announce routes of a certain source type to BGP, for example, DvR, direct, static, Routing Information Protocol (RIP), and Open Shortest Path First (OSPF). If you do not configure a route policy, then the switch uses the default action based on metric, metric type, and subnet. Use a route policy to perform detailed redistribution.

Before you begin

- If required, configure a route policy.
- When you configure BGP on a specific VRF instance, the VRF must have an RP trigger of BGP. Not all parameters are configurable on non-default VRFs.
- Before you redistribute DvR host routes to BGP, ensure that you disable BGP aggregation and BGP auto-summarization of networks. Disabling these settings ensures that all the DvR host routes are advertised into BGP correctly, and are not summarized.

Procedure

1. In the navigation pane, open the following folders: **Configuration > IP**.

2. Click **BGP**.
3. Click the **Redistribute** tab.
4. Click **Insert**.
5. Configure the source protocol.
6. If required, choose a route policy.
7. Configure the metric to apply to redistributed routes.
8. Enable the redistribution instance.
9. Click **Insert**.

Redistribute field descriptions

Use the data in the following table to use the **Redistribute** tab.

Name	Description
DstVrflid	Specifies the destination VRF instance (read-only).
Protocol	Specifies the protocols that receive the redistributed routes.
SrcVrflid	Specifies the source VRF instance (read-only).
RouteSource	Specifies the source protocol for the route redistribution entry.
Enable	Enables (or disables) a BGP redistribute entry for a specified source type.
RoutePolicy	Configures the route policy to use for the detailed redistribution of external routes from a specified source into the BGP domain.
Metric	Configures the metric for the redistributed route. The value can be a range between 0–65535. The default value is 0. Use a value that is consistent with the destination protocol.

Configuring redistribution to BGPv6

Configure redistribute entries for BGPv6 to announce routes of a certain source type to BGPv6, for example, DvR, direct, static, Routing Information Protocol (RIP), and Open Shortest Path First (OSPF). If you do not configure a route policy, then the switch uses the default action based on metric, metric type, and subnet. Use a route policy to perform detailed redistribution.

Before you begin

- If required, configure a route policy.
- Before you redistribute DvR host routes to BGPv6, ensure that you disable BGPv6 aggregation and BGPv6 autosummarization of networks. Disabling these settings ensures that all the DvR host routes are advertised into BGPv6 correctly, and are not summarized.

Procedure

1. In the navigation pane, expand the following folders: **Configuration > IPv6**.
2. Click **BGP+**.
3. Click the **Redistribute** tab.
4. Click **Insert**.
5. Configure the source protocol.
6. If required, choose a route policy.
7. Configure the metric to apply to redistributed routes.
8. Enable the redistribution instance.
9. Click **Insert**.

Redistribute field descriptions

Use the data in the following table to use the Redistribute tab.

Name	Description
DstVrflid	Specifies the destination VRF instance (read-only).
Protocol  Note: This field does not appear on all hardware platforms.	Specifies the protocols that receive the redistributed routes.
SrcVrflid	Specifies the source VRF instance (read-only).
RouteSource	Specifies the source protocol for the route redistribution entry.
Enable	Enables (or disables) a BGPv6 redistribute entry for a specified source type.
RoutePolicy	Configures the route policy to use for the detailed redistribution of external routes from a specified source into the BGPv6 domain.
Metric	Configures the metric for the redistributed route. The default value is 0. Use a value that is consistent with the destination protocol.

Viewing BGP+ or BGPv6 route summary information

You can display current IPv6 BGP+ route information.

Procedure

1. In the navigation tree, expand the following folders: **Configuration > IPv6**.

Use the **IP** folder in the navigation tree to view the IPv4 route summary.

2. Click **BGP+**.
3. Click the **Bgp Route Summary** tab to view the BGP route summary information

Bgp Route Summary field descriptions

Use the data in the following table to use the **Bgp Route Summary** tab.

Name	Description
Prefix	Specifies the IP address prefix in the Network Layer Reachability Information (NLRI) field. This is an IP address that contains the prefix with a length specified by IpAddrPrefixLen. Any bits beyond the length specified by IpAddrPrefixLen are set to zero.
PrefixLen	Specifies the length, in bits, of the IP address prefix in the NLRI field.
LocalAddr	The local address of this entry's BGP connection.
RemoteAddr	Specifies the IP address of the peer from which path information was learned.

Viewing BGP route summary

Display BGP route summary.

Procedure

1. In the navigation pane, open the following folders: **Configuration > IP**.
2. Click **BGP**.
3. Click the **Bgp Route Summary** tab.

Bgp Route Summary field descriptions

Use the data in the following table to use the Bgp Route Summary tab.

Name	Description
Prefix	Configures the IP address of the route.
PrefixLen	Specifies the IP address and the mask length (the length can be 0–32).

Table continues...

Name	Description
LocalAddr	Specifies the local IP address of the entered BGP route.
RemoteAddr	Specifies the remote IP address of the entered BGP route.

Configuring an AS path list

Configure an AS path list to restrict the routing information a router learns or advertises to and from a neighbor. The AS path list acts as a filter that matches AS paths.

Before you begin

- Change the VRF instance as required to configure BGP on a specific VRF instance. The VRF must have an RP trigger of BGP. Not all parameters are configurable on non-default VRFs.

Procedure

1. In the navigation pane, open the following folders: **Configuration > IP**.
2. Click **Policy**.
3. Click the **As Path List** tab.
4. Click **Insert**.
5. Enter the appropriate information for your configuration.
6. Click **Insert**.

As Path List field descriptions

Use the data in the following table to use the As Path List tab.

Name	Description
Id	Specifies the AS path list.
MemberId	Specifies the AS path access list member ID.
Mode	Specifies the action to take if the system selects a policy for a specific route. Select permit (allow the route) or deny (ignore the route).
AsRegularExpression	Specifies the expression to use for the AS path.

Configuring a community access list

Configure community lists to specify permitted routes by using their BGP community. This list acts as a filter that matches communities or AS numbers.

Before you begin

- Change the VRF instance as required to configure BGP on a specific VRF instance. The VRF must have an RP trigger of BGP. Not all parameters are configurable on non-default VRFs.

Procedure

1. In the navigation pane, open the following folders: **Configuration > IP**.
2. Click **Policy**.
3. Click the **Community List** tab.
4. Click **Insert**.
5. Configure the list as required.
6. Click **Insert**.

Community List field descriptions

Use the data in the following table to use the **Community List** tab.

Name	Description
Id	Specifies the community list. The range is 0–1024.
MemberId	Specifies the community list member ID. The range is 0–65535.
Mode	Specifies the action to take if the system selects a policy for a specific route. Select permit (allow the route) or deny (ignore the route).
Community	Specifies the community access list community string.

Chapter 7: BGP Configuration Examples

This chapter shows configuration examples for BGP deployment options.

IPv6 Tunnel Configurations for BGP+

You must configure an IPv6 tunnel and static routes at BGP+ peers when you use BGP+.

When BGP+ peers advertise route information, they use Update messages to advertise route information. And, when route information is encapsulated in Update messages, BGP+ peers convert their own IPv4 peer addresses to IPv4-mapped IPv6 addresses and insert them into the next-hop field in the Update message.

When the BGP+ software module receives Update messages, it adds route information to the IPv6 Routing Manager (RTM). These RTM routes contain next-hop addresses from the BGP peer that the route was learned from. The next-hop addresses are represented as IPv4-mapped IPv6 addresses.

But, because the IPv6 RTM cannot correlate the IPv4-mapped IPv6 address to a specific outgoing interface, you must create a manually-configured static route to make the link between the BGP peer and the IPv6 tunnel interface so that traffic can reach networks advertised by the peer.

Following is one way to express a static route in an IPv6-configured tunnel for BGP+:

```
ipv6 route 0:0:0:0:0:ffff:192.0.2.0/24 cost 1 tunnel 10
```

It is recommended that the IPv6-configured tunnel endpoint and the BGP peer reside on the same switch.

If the IPv6 tunnel endpoint and the BGP peer must reside on different switches you can terminate the tunnel on a different switch, but you must consider the following:

- Because the IPv6 tunnel endpoint does not reside on the same switch as the BGP peer, the BGP device cannot use the tunnel as the outgoing interface. That is, to reach the IPv6-configured tunnel endpoint, if the BGP peer resides on a different switch from the IPv6 tunnel endpoint, the next-hop for the manually-configured IPv4-mapped IPv6 static route is the native IPv6 interface next-hop address.
- The node where the tunnel terminates must contain all of the information needed to route the packets between the remote IPv6 network clouds.

*** Note:**

In order for the tunnel endpoint switch to be aware of all of the necessary IPv6 routes, you may need to redistribute the BGP routes into OSPFv3.

IPv4-mapped IPv6 address

IPv4-mapped IPv6 addresses are IPv4 addresses that the system has mapped into the IPv6 address space.

The system uses these IPv4-mapped IPv6 addresses for devices that are only IPv4-capable.

These IPv4-mapped address have the first 80 bits set to zeros, followed by the next 16 bits set to ones, and the last 32 bits have IPv4 addresses.

When converted to an IPv4-mapped IPv6 address, an IPv4 device address of 192.0.2.1 would be represented as one of the following:

- 0:0:0:0:FFFF:192.0.2.1
- ::FFFF:192.0.2.1

The following figure illustrates the components in an IPv4-mapped IPv6 address.

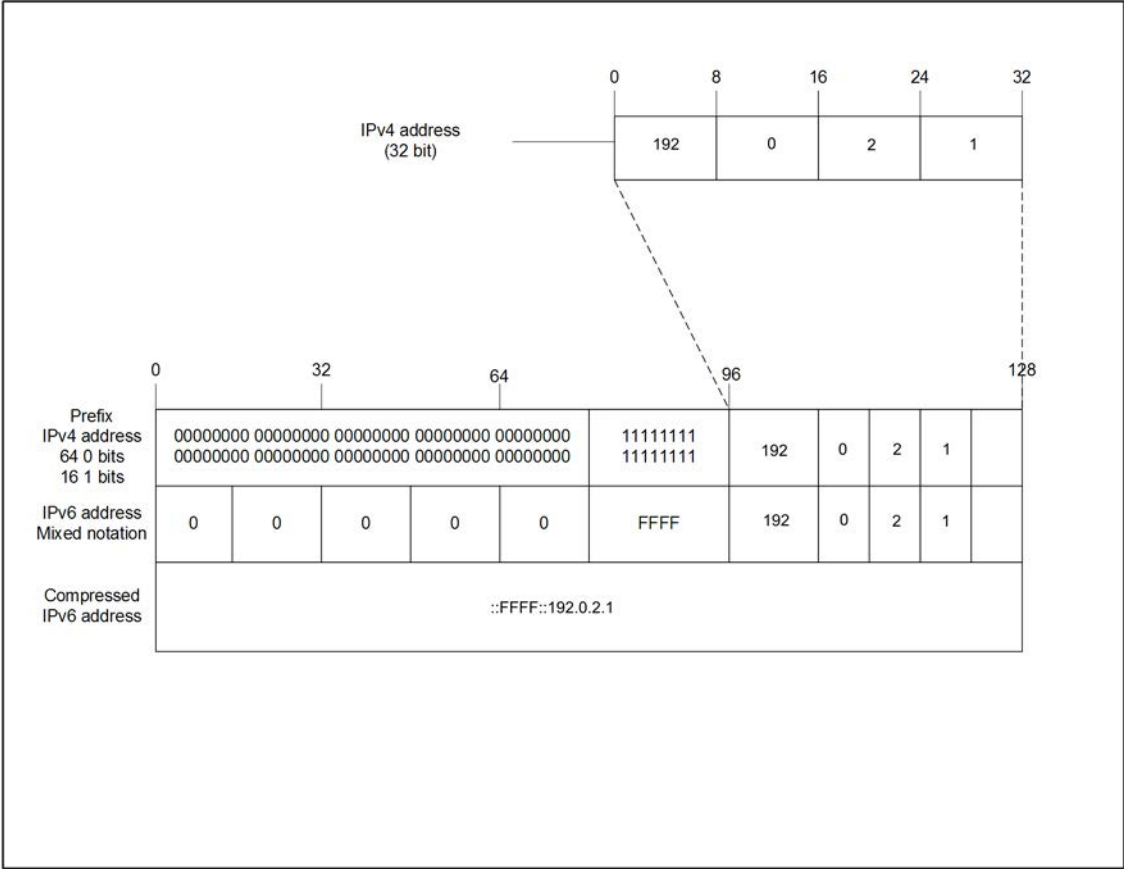


Figure 23: IPv4-mapped IPv6 address components

eBGP+ peership between two switches with IPv6 Tunneling

The following figure shows a sample network that contains eBGP+ peers using IPv6 tunneling.

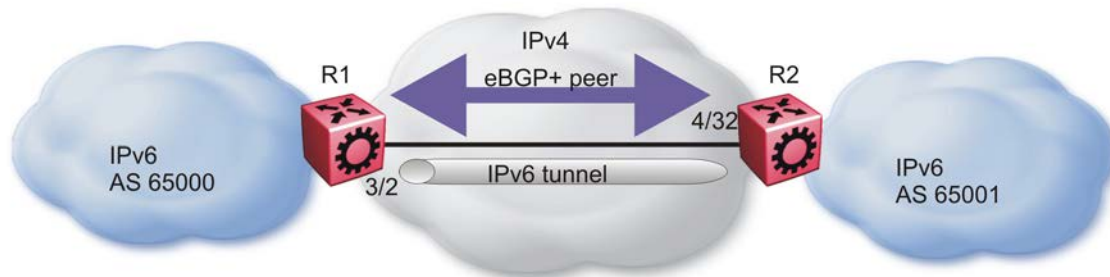


Figure 24: eBGP+ peers with IPv6 tunneling

The configuration in the figure, *eBGP+ peers with IPv6 tunneling*, assumes that the BGP peer IP address is the next hop.

When you configure the static route for the BGP+ tunnel, you must designate the BGP peer IP address as the next hop in most cases.

You can configure multiple static routes, using the same tunnel, but you must ensure reachability when you create the static routes.

R1 configuration

```
interface GigabitEthernet 3/2
brouter port 3/2 vlan 2090 subnet 192.0.2.1/255.255.255.0 mac-offset 2
exit
# BGP CONFIGURATION - GlobalRouter
#

router bgp as-dot enable
router bgp 65000 enable
router bgp
neighbor "192.0.2.2"
neighbor 192.0.2.2 remote-as 65001
neighbor 192.0.2.2 address-family ipv6
neighbor 192.0.2.2 enable
exit
# IPV6 CONFIGURATION
#

ipv6 forwarding

# IPV6 TUNNEL CONFIGURATION
#

ipv6 tunnel 10 source 192.0.2.1 address 2001:DB8::/32 destination 200.1.
1.2

#
# IPV6 STATIC ROUTE CONFIGURATION
#

ipv6 route 0:0:0:0:0:ffff:192.0.2.1 cost 1 tunnel 10
#
```

R2 configuration

```

interface GigabitEthernet 4/32
brouter port 4/32 vlan 2090 subnet 192.0.2.2/255.255.255.0 mac-offset
2
exit
# BGP CONFIGURATION - GlobalRouter
#

router bgp as-dot enable
router bgp 65000 enable
router bgp
neighbor "192.0.2.1"
neighbor 192.0.2.1 remote-as 65000
neighbor 192.0.2.1 address-family ipv6
neighbor 192.0.2.1 enable
exit
# IPV6 CONFIGURATION
#

ipv6 forwarding

# IPV6 TUNNEL CONFIGURATION
#

ipv6 tunnel 10 source 192.0.2.2 address 2001:DB8::/32 destination
192.0.2.1

#
# IPV6 STATIC ROUTE CONFIGURATION
#

ipv6 route 0:0:0:0:0:ffff:192.0.2.1 cost 1 tunnel 10
#

```

iBGP+ peership on CLIP between two switches with IPv6 Tunneling

The following figure shows a sample network that contains iBGP+ peers using IPv6 tunneling.

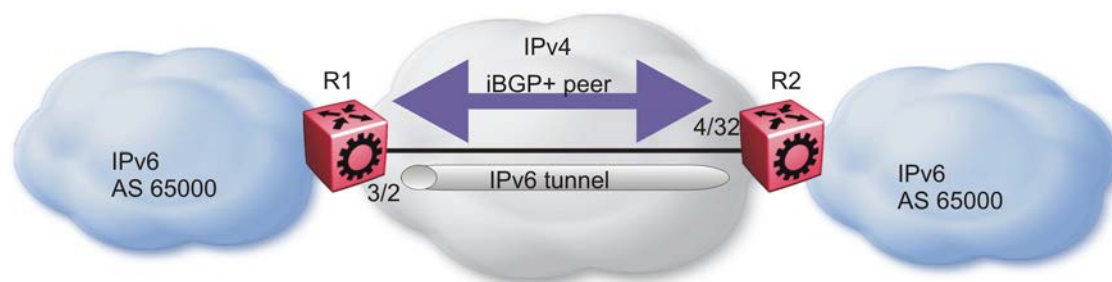


Figure 25: iBGP+ peers on CLIP interfaces with IPv6 tunneling

You must enable OSPF on the interface and globally as well.

If you cannot enable OSPF, you must configure static routes to provide reachability to the BGP+ peer.

The static route must point to the next hop for the routes to be installed in the IPv6 RTM.

The next hop must be the BGP peer IP address.

The IPv4 interfaces do not need to connect directly, but the routing table on each switch must include the IPv4 interface of the other switch.

iBGP between the CLIP interfaces needs to run OSPF as a routing protocol so that the BGP neighbor can remain reachable.

eBGP connections cannot use a CLIP interface as an end point.

R1 configuration

```
interface GigabitEthernet 3/2
brouter port 3/2 vlan 2090 subnet 192.0.2.1/255.255.255.0 mac-offset
2
exit
# OSPF CONFIGURATION - GlobalRouter
#

router ospf enable

# OSPF PORT CONFIGURATION
#

interface gigabitethernet 3/2
ip ospf enable
exit

# CIRCUITLESS IP INTERFACE CONFIGURATION - GlobalRouter
#

interface loopback 1
ip address 1 1.1.1.1/255.255.255.255
ip ospf 1

# BGP CONFIGURATION - GlobalRouter
#

router bgp
no synchronization
exit
router bgp as-dot enable
router bgp 65000 enable
router bgp
neighbor "2.2.2.2"
neighbor 2.2.2.2 remote-as 65000
neighbor 2.2.2.2 next-hop-self
neighbor 2.2.2.2 update-source 1.1.1.1
neighbor 2.2.2.2 address-family ipv6
neighbor 2.2.2.2 enable
exit
# IPV6 CONFIGURATION
#

ipv6 forwarding

# IPV6 TUNNEL CONFIGURATION
#

ipv6 tunnel 10 source 192.0.2.1 address 2001:DB8::/32 destination
192.0.2.2

#
# IPV6 STATIC ROUTE CONFIGURATION
```

```
#  
ipv6 route 0:0:0:0:0:ffff:2.2.2.2/128 cost 1 tunnel 10  
#
```

R2 configuration

```
interface GigabitEthernet 4/32  
brouter port 4/32 vlan 2090 subnet 192.0.2.2/255.255.255.0 mac-offset  
2  
exit  
# OSPF CONFIGURATION - GlobalRouter  
#  
router ospf enable  
  
# OSPF PORT CONFIGURATION  
#  
interface gigabitethernet 4/32  
ip ospf enable  
exit  
  
# CIRCUITLESS IP INTERFACE CONFIGURATION - GlobalRouter  
#  
interface loopback 1  
ip address 1 2.2.2.2/255.255.255.255  
ip ospf 1  
  
# BGP CONFIGURATION - GlobalRouter  
#  
router bgp  
no synchronization  
exit  
router bgp as-dot enable  
router bgp 65000 enable  
router bgp  
neighbor "1.1.1.1"  
neighbor 1.1.1.1 remote-as 65000  
neighbor 1.1.1.1 next-hop-self  
neighbor 1.1.1.1 update-source 2.2.2.2  
neighbor 1.1.1.1 address-family ipv6  
neighbor 1.1.1.1 enable  
exit  
# IPV6 CONFIGURATION  
#  
ipv6 forwarding  
  
# IPV6 TUNNEL CONFIGURATION  
#  
ipv6 tunnel 10 source 192.0.2.2 address 2001:DB8::/32 destination  
192.0.2.1  
  
#  
# IPV6 STATIC ROUTE CONFIGURATION  
#  
ipv6 route 0:0:0:0:0:ffff:1.1.1.1/128 cost 1 tunnel 10  
#
```

Native IPv6 eBGP peership between two switches on VRF

The following figure shows a sample network that contains native IPv6 eBGP peers on VRFs.



Figure 26: Native IPv6 eBGP peers on VRFs

Configure the local AS first on GRT (and it is inherited by all VRFs), and then enable BGP on GRT/VRF.

You must configure the **address-family ipv6** option for IPv6 peers, otherwise, peer-ship is formed, but no routing updates between them will take place.

You must configure the **ebgp-multihop** option for the given eBGP peer that is not on one of local subnets (remote peers), otherwise, peer-ship will not be formed.

* Note:

The switch does not accept any configuration command for BGP in router-vrf configuration mode unless a BGP instance associated to the VRF context is created. You can use `ip bgp` command in router-vrf configuration mode to create a BGP instance on VRF.

R1 configuration

```
#
# VRF CONFIGURATION
#
ip vrf vrf1 vrfid 1
router vrf vrf1
exit
ip vrf vrf2 vrfid 2
router vrf vrf2
exit
#

# PORT CONFIGURATION - PHASE I
#
interface GigabitEthernet 1/1
encapsulation dot1q
exit

#
# VLAN CONFIGURATION
#
vlan members remove 1 1/1,1/46
vlan create 100 type port-mstprstp 0
```

BGP Configuration Examples

```
vlan members 100 1/1 portmember
interface Vlan 100
ip address 100.1.1.1 255.255.255.0 1
ipv6 interface mac-offset 1
ipv6 interface enable
ipv6 interface address 2001:0:100:0:0:0:1/64
exit
vlan create 101 type port-mstprstp 0
vlan members 101 1/1 portmember
interface Vlan 101
vrf vrf1
ip address 101.1.1.1 255.255.255.0 2
ipv6 interface mac-offset 2
ipv6 interface enable
ipv6 interface address 2001:0:101:0:0:0:1/64
exit
vlan create 102 type port-mstprstp 0
vlan members 102 1/1 portmember
interface Vlan 102
vrf vrf2
ip address 102.1.1.1 255.255.255.0 3
ipv6 interface mac-offset 3
ipv6 interface enable
ipv6 interface address 2001:0:102:0:0:0:1/64
ipv6 forwarding

exit

#
# PORT CONFIGURATION - PHASE II
#

interface GigabitEthernet 1/1
default-vlan-id 100
no shutdown
exit

#
# CIRCUITLESS IPV6 INTERFACE CONFIGURATION - GlobalRouter
#

interface loopback 1
ipv6 interface address 1:1:1:1:0:0:0:1/128

exit
#
# CIRCUITLESS IPV6 INTERFACE CONFIGURATION - VRF
#

interface loopback 2
ipv6 interface address 11:1:1:1:0:0:0:1/128 vrf vrf1
exit
interface loopback 3
ipv6 interface address 12:1:1:1:0:0:0:1/128 vrf vrf2
exit

#
# BGP CONFIGURATION - GlobalRouter
#

router bgp
no synchronization
exit
router bgp 1000 enable
router bgp
```

```

network 1:1:1:1:0:0:0:1/128 metric 100000
neighbor "2001:0:100:0:0:0:0:2"
neighbor 2001:0:100:0:0:0:0:2 remote-as 10000
neighbor 2001:0:100:0:0:0:0:2 next-hop-self
neighbor 2001:0:100:0:0:0:0:2 ebgp-multihop
neighbor 2001:0:100:0:0:0:0:2 address-family ipv6
neighbor 2001:0:100:0:0:0:0:2 update-source 2001:0:100:0:0:0:1
neighbor 2001:0:100:0:0:0:0:2 enable
exit#
# BGP CONFIGURATION - VRF
#
router vrf vrf1
ip bgp
no ip bgp synchronization
ip bgp enable
ip bgp network 11:1:1:1:0:0:0:1/128 metric 100000
ip bgp neighbor "2001:0:101:0:0:0:0:2"
ip bgp neighbor 2001:0:101:0:0:0:0:2 remote-as 10000
ip bgp neighbor 2001:0:101:0:0:0:0:2 next-hop-self
ip bgp neighbor 2001:0:101:0:0:0:0:2 ebgp-multihop
ip bgp neighbor 2001:0:101:0:0:0:0:2 address-family ipv6
ip bgp neighbor 2001:0:101:0:0:0:0:2 update-source 2001:0:101:0:0:0:1
ip bgp neighbor 2001:0:101:0:0:0:0:2 enable
exit
router vrf vrf2
ip bgp
no ip bgp synchronization
ip bgp enable
ip bgp network 12:1:1:1:0:0:0:1/128 metric 100000
ip bgp neighbor "2001:0:102:0:0:0:0:2"
ip bgp neighbor 2001:0:102:0:0:0:0:2 remote-as 10000
ip bgp neighbor 2001:0:102:0:0:0:0:2 next-hop-self
ip bgp neighbor 2001:0:102:0:0:0:0:2 ebgp-multihop
ip bgp neighbor 2001:0:101:0:0:0:0:2 address-family ipv6
ip bgp neighbor 2001:0:102:0:0:0:0:2 update-source 2001:0:102:0:0:0:1
ip bgp neighbor 2001:0:102:0:0:0:0:2 enable
exit

```

R2 configuration

```

#
# VRF CONFIGURATION
#

ip vrf vrf1 vrfid 1
router vrf vrf1
exit
ip vrf vrf2 vrfid 2
router vrf vrf2
exit

#
# PORT CONFIGURATION - PHASE I
#

interface GigabitEthernet 1/1
encapsulation dot1q

exit

#
# VLAN CONFIGURATION
#

vlan members remove 1 1/1

```

BGP Configuration Examples

```
vlan create 100 type port-mstprstp 0
vlan members 100 1/1 portmember
interface Vlan 100

ip address 100.1.1.2 255.255.255.0 1
ipv6 interface mac-offset 1
ipv6 interface enable
ipv6 interface address 2001:0:100:0:0:0:2/64
ipv6 forwarding

exit
vlan create 101 type port-mstprstp 0
vlan members 101 1/1 portmember
interface Vlan 101
vrf vrf1
ip address 101.1.1.2 255.255.255.0 2
ipv6 interface mac-offset 2
ipv6 interface enable
ipv6 interface address 2001:0:101:0:0:0:2/64
ipv6 forwarding

exit
vlan create 102 type port-mstprstp 0
vlan members 102 1/1 portmember
interface Vlan 102
vrf vrf2
ip address 102.1.1.2 255.255.255.0 3

ipv6 interface mac-offset 3
ipv6 interface enable
ipv6 interface address 2001:0:102:0:0:0:2/64
ipv6 forwarding

exit

#
# PORT CONFIGURATION - PHASE II
#

interface GigabitEthernet 1/1
default-vlan-id 100
no shutdownexit

#
# CIRCUITLESS IPV6 INTERFACE CONFIGURATION - GlobalRouter
#

interface loopback 1
ipv6 interface address 2:2:2:2:0:0:0:2/128
exit

#
# CIRCUITLESS IPV6 INTERFACE CONFIGURATION - VRF
#

interface loopback 2
ipv6 interface address 21:2:2:2:0:0:0:2/128 vrf vrf1
exit
interface loopback 3
ipv6 interface address 22:2:2:2:0:0:0:2/128 vrf vrf2
exit

#
```

```

# BGP CONFIGURATION - GlobalRouter
#

router bgp
no synchronization
exit
router bgp 10000 enable
router bgp
neighbor "2001:0:100:0:0:0:0:1"
neighbor 2001:0:100:0:0:0:0:1 remote-as 1000
neighbor 2001:0:100:0:0:0:0:1 next-hop-self
neighbor 2001:0:100:0:0:0:0:1 ebgp-multihop
neighbor 2001:0:100:0:0:0:0:1 address-family ipv6
neighbor 2001:0:100:0:0:0:0:1 update-source 2001:0:100:0:0:0:0:2
neighbor 2001:0:100:0:0:0:0:1 enableexit

#
# BGP CONFIGURATION - VRF
#

router vrf vrf1
ip bgp
no ip bgp synchronization
ip bgp enable
ip bgp neighbor "2001:0:101:0:0:0:0:1"
ip bgp neighbor 2001:0:101:0:0:0:0:1 remote-as 1000
ip bgp neighbor 2001:0:101:0:0:0:0:1 next-hop-self
ip bgp neighbor 2001:0:101:0:0:0:0:1 ebgp-multihop
ip bgp neighbor 2001:0:101:0:0:0:0:1 address-family ipv6
ip bgp neighbor 2001:0:101:0:0:0:0:1 update-source 2001:0:101:0:0:0:0:2
ip bgp neighbor 2001:0:101:0:0:0:0:1 enable
exit
router vrf vrf2
ip bgp
no ip bgp synchronization
ip bgp enable
ip bgp neighbor "2001:0:102:0:0:0:0:1"
ip bgp neighbor 2001:0:102:0:0:0:0:1 remote-as 1000
ip bgp neighbor 2001:0:102:0:0:0:0:1 next-hop-self
ip bgp neighbor 2001:0:102:0:0:0:0:1 ebgp-multihop
ip bgp neighbor 2001:0:102:0:0:0:0:1 address-family ipv6
ip bgp neighbor 2001:0:102:0:0:0:0:1 update-source 2001:0:102:0:0:0:0:2
ip bgp neighbor 2001:0:102:0:0:0:0:1 enable
exit

```

Glossary

4-byte AS	4-byte Autonomous System (AS) numbers is the solution to the soon depleting 2-byte AS numbers. It provides a theoretical 4,294,967,296 unique AS numbers in BGP. 4-byte AS numbers are backward compatible with 2-byte AS numbers.
AS confederation	A single logical autonomous system (AS) that comprises of multiple sub-autonomous systems to ensure scalability.
AS_TRANS	RFC4893 defines a 4-octet Autonomous System number 23456 to facilitate backward compatibility. This 23456 AS number is also known as AS_TRANS (AS_TRANS=23456).
attribute	A unit of data BGP used to describe the prefixes, such as AS-PATH, LOCAL-PREF, NEXT-HOP.
Autonomous System (AS)	A set of routers under a single technical administration, using a single IGP and common metrics to route packets within the Autonomous System, and using an EGP to route packets to other Autonomous Systems.
Autonomous System Number (ASN)	A two-byte number that is used to identify a specific AS.
Border Gateway Protocol (BGP)	An inter-domain routing protocol that provides loop-free inter-domain routing between Autonomous Systems (AS) or within an AS.
Border Gateway Protocol neighbor	Border Gateway Protocol routers that have interfaces to a common network.
Border Gateway Protocol peer	A relationship that is formed between two routers that open a TCP connection to each other for the purpose of exchanging routing information.
Border Gateway Protocol session	An active connection between two routers running BGP.
Border Gateway Protocol speaker	An entity within a BGP router that is used to communicate with other BGP speakers by establishing a peer-to-peer session.
Circuitless IP (CLIP)	A CLIP is often called a loopback and is a virtual interface that does not map to any physical interface.

classless interdomain routing (CIDR)	The protocol defined in RFCs 1517 and 1518 for using subnetwork masks, other than the defaults for IP address classes.
cluster	One or more route reflectors and their associated clients that form a relationship where the designated route reflectors provide route reflection for their clients, as well as nonclient peers.
community	A BGP attribute that contains a list of 32-bit values used to identify a route as belonging to a category of routes. All of the routes in the category are treated equally by routing policies.
dampen	Indicates that routes which exhibit instability are not advertised until the routes become stable for a minimum time period.
equal cost multipath (ECMP)	Distributes routing traffic among multiple equal-cost routes.
External BGP (eBGP)	A Border Gateway Protocol (BGP) used by routers that exchange information between two BGP speakers in different Autonomous Systems.
Interior BGP (iBGP)	Routers that use the Border Gateway Protocol (BGP) within an Autonomous System. The router redistributes BGP information to Interior Gateway Protocols (IGPs) that run in the autonomous path.
Interior Gateway Protocol (IGP)	Distributes routing information between routers that belong to a single Autonomous System (AS).
Internet Assigned Numbers Authority (IANA)	The central registry for various assigned numbers, for example, Internet protocol parameters (such as port, protocol, and enterprise numbers), options, codes, and types.
load balancing	The practice of splitting communication into two (or more) routes or servers.
Message Digest 5 (MD5)	A one-way hash function that creates a message digest for digital signatures.
multihomed AS	An autonomous system that has multiple connections to one or more autonomous systems and does not carry transit traffic.
next hop	The next hop to which a packet can be sent to advance the packet to the destination.
prefix	A group of contiguous bits, from 0 to 32 bits in length, that defines a set of addresses.
route reflector	A BGP speaker that advertises routes learned from its route reflector clients to other iBGP neighbors.

route reflector client	A BGP speaker that advertises its learned routes to a route reflector for readvertisement of its routes to the rest of the AS.
routing policy	A form of routing that is influenced by factors other than the default algorithmically best route, such as the shortest or quickest path.
transit AS	An autonomous system (AS) that has multiple connections to one or more autonomous systems and is used (with certain policy restrictions) to carry both transit and local traffic.
well-known attribute	A BGP attribute that is required to be known by all BGP implementations.