



Border Gateway Protocol (BGP)

BGP is used to exchange routing information for the Internet and is the protocol used between Internet service providers (ISP). If your system is a gateway to the service provider network, you might need to implement BGP. You can configure one BGP process on the device, for a single autonomous system.

- [About BGP, on page 1](#)
- [Configure BGP, on page 4](#)
- [Monitoring BGP, on page 24](#)

About BGP

BGP is an inter and intra autonomous system routing protocol. An autonomous system is a network or group of networks under a common administration and with common routing policies. BGP is used to exchange routing information for the Internet and is the protocol used between Internet service providers (ISP).

Routing Table Changes

BGP neighbors exchange full routing information when the TCP connection between neighbors is first established. When changes to the routing table are detected, the BGP routers send to their neighbors only those routes that have changed. BGP routers do not send periodic routing updates, and BGP routing updates advertise only the optimal path to a destination network.



Note AS loop detection is done by scanning the full AS path (as specified in the AS_PATH attribute), and checking that the AS number of the local system does not appear in the AS path. By default, EBGP advertises the learned routes to the same peer to prevent additional CPU cycles on the device in performing loop checks and to avoid delays in the existing outgoing update tasks.

Routes learned via BGP have properties that are used to determine the best route to a destination, when multiple paths exist to a particular destination. These properties are referred to as BGP attributes and are used in the route selection process:

- **Weight**—This is a Cisco-defined attribute that is local to a router. The weight attribute is not advertised to neighboring routers. If the router learns about more than one route to the same destination, the route with the highest weight is preferred.

- **Local preference**—The local preference attribute is used to select an exit point from the local AS. Unlike the weight attribute, the local preference attribute is propagated throughout the local AS. If there are multiple exit points from the AS, the exit point with the highest local preference attribute is used as an exit point for a specific route.
- **Multi-exit discriminator**—The multi-exit discriminator (MED) or metric attribute is used as a suggestion to an external AS regarding the preferred route into the AS that is advertising the metric. It is referred to as a suggestion because the external AS that is receiving the MEDs may also be using other BGP attributes for route selection. The route with the lower MED metric is preferred.
- **Origin**—The origin attribute indicates how BGP learned about a particular route. The origin attribute can have one of three possible values and is used in route selection.
 - **IGP**—The route is interior to the originating AS. This value is set when the network router configuration command is used to inject the route into BGP.
 - **EGP**—The route is learned via the Exterior Border Gateway Protocol (EBGP).
 - **Incomplete**—The origin of the route is unknown or learned in some other way. An origin of incomplete occurs when a route is redistributed into BGP.
- **AS_path**—When a route advertisement passes through an autonomous system, the AS number is added to an ordered list of AS numbers that the route advertisement has traversed. Only the route with the shortest AS_path list is installed in the IP routing table.
- **Next hop**—The EBGP next-hop attribute is the IP address that is used to reach the advertising router. For EBGP peers, the next-hop address is the IP address of the connection between the peers. For IBGP, the EBGP next-hop address is carried into the local AS. However, when the next hop is in the same subnet as the peering address of the eBGP peer, the next hop is not modified. This behavior is referred to as the third party next hop.
- **Community**—The community attribute provides a way of grouping destinations, called communities, to which routing decisions (such as acceptance, preference, and redistribution) can be applied. Route maps are used to set the community attribute. The predefined community attributes are as follows:
 - **no-export**—Do not advertise this route to EBGP peers.
 - **no-advertise**—Do not advertise this route to any peer.
 - **internet**—Advertise this route to the Internet community; all routers in the network belong to it.

When to Use BGP

Customer networks, such as universities and corporations, usually employ an Interior Gateway Protocol (IGP) such as OSPF for the exchange of routing information within their networks. Customers connect to ISPs, and ISPs use BGP to exchange customer and ISP routes. When BGP is used between autonomous systems (AS), the protocol is referred to as External BGP (EBGP). If a service provider is using BGP to exchange routes within an AS, then the protocol is referred to as Interior BGP (IBGP).

BGP can also be used for carrying routing information for IPv6 prefix over IPv6 networks.

BGP Path Selection

BGP may receive multiple advertisements for the same route from different sources. BGP selects only one path as the best path. When this path is selected, BGP puts the selected path in the IP routing table and propagates the path to its neighbors. BGP uses the following criteria, in the order presented, to select a path for a destination:

- If the path specifies a next hop that is inaccessible, drop the update.
- Prefer the path with the largest weight.
- If the weights are the same, prefer the path with the largest local preference.
- If the local preferences are the same, prefer the path that was originated by BGP running on this router.
- If no route was originated, prefer the route that has the shortest AS_path.
- If all paths have the same AS_path length, prefer the path with the lowest origin type (where IGP is lower than EGP, and EGP is lower than incomplete).
- If the origin codes are the same, prefer the path with the lowest MED attribute.
- If the paths have the same MED, prefer the external path over the internal path.
- If the paths are still the same, prefer the path through the closest IGP neighbor.
- Determine if multiple paths require installation in the routing table for [BGP Multipath, on page 3](#).
- If both paths are external, prefer the path that was received first (the oldest one).
- Prefer the path with the lowest IP address, as specified by the BGP router ID.
- If the originator or router ID is the same for multiple paths, prefer the path with the minimum cluster list length.
- Prefer the path that comes from the lowest neighbor address.

BGP Multipath

BGP Multipath allows installation into the IP routing table of multiple equal-cost BGP paths to the same destination prefix. Traffic to the destination prefix is then shared across all installed paths.

These paths are installed in the table together with the best path for load-sharing. BGP Multipath does not affect best-path selection. For example, a router still designates one of the paths as the best path, according to the algorithm, and advertises this best path to its BGP peers.

In order to be candidates for multipath, paths to the same destination need to have these characteristics equal to the best-path characteristics:

- Weight
- Local preference
- AS-PATH length
- Origin code
- Multi Exit Discriminator (MED)

- One of these:
 - Neighboring AS or sub-AS (before the addition of the BGP Multipaths)
 - AS-PATH (after the addition of the BGP Multipaths)

Some BGP Multipath features put additional requirements on multipath candidates:

- The path should be learned from an external or confederation-external neighbor (eBGP).
- The IGP metric to the BGP next hop should be equal to the best-path IGP metric.

These are the additional requirements for internal BGP (iBGP) multipath candidates:

- The path should be learned from an internal neighbor (iBGP).
- The IGP metric to the BGP next hop should be equal to the best-path IGP metric, unless the router is configured for unequal-cost iBGP multipath.

BGP inserts up to n most recently received paths from multipath candidates into the IP routing table, where n is the number of routes to install to the routing table, as specified when you configure BGP Multipath. The default value, when multipath is disabled, is 1.

For unequal-cost load balancing, you can also use BGP Link Bandwidth.



Note The equivalent next-hop-self is performed on the best path that is selected among eBGP multipaths before it is forwarded to internal peers.

Configure BGP

The following topics explain how to configure BGP.

Configure BGP Global Settings

If you configure BGP, the global settings apply across all virtual routers, if you use virtual routers. There are additional BGP settings you configure to define the BGP process. When using virtual routers, you can create a separate BGP process for each virtual router.

Before you begin

After you create the BGP Global Settings object, you can delete it if you no longer need it. Simply edit the object following this procedure, but click the **Delete BGP Global Settings Object** button at the bottom of the dialog box.

Procedure

-
- Step 1** Click **Device**, then click the **Routing** summary.
 - Step 2** On the main routing or virtual routers page, click the **BGP Global Settings** button.

If you are viewing a virtual router, you must return to the main list of virtual routers.

- Step 3** If you have not yet configured the BGP Global Settings object, click **Create BGP Global Settings Object**.
- Step 4** (Optional.) You can change the object name or enter a description for the object. The default object name is BgpGeneralSettings.
- Step 5** Configure at least the following basic settings:
- **router bgp as-number**. Click *as-number* and enter the autonomous system (AS) number for the BGP process. The AS number can be from 1 to 4294967295 or from 1.0 to 65535.65535. The AS number is a uniquely assigned value that identifies each network on the Internet. The system supports asplain and asdot notation as defined in RFC 5396.
 - **log-neighbor-changes state**. Click *state* and select enable or disable. When enabled, which is recommended, BGP neighbor changes (up or down) and resets are logged. This helps in troubleshooting network connectivity problems and measuring network stability.
 - **transport path-mtu-discovery state**. Click *state* and select enable or disable. When enabled, which is recommended, the system determines the maximum transmission unit (MTU) size on the network path between two IP hosts, and then takes advantage of the highest-MTU path. This avoids IP fragmentation.
 - **fast-external-fallover state**. Click *state* and select enable or disable. When enabled, which is recommended, the system uses fast external fallover for BGP peering sessions with directly connected external peers. The session is immediately reset if the link goes down. If you disable BGP fast external fallover, the BGP routing process will wait until the default hold timer expires (3 keepalives) to reset the peering session.
 - **enforce-first-as state**. Click *state* and select enable or disable. When enabled, which is recommended, the system denies incoming updates received from eBGP peers that do not list their autonomous system number as the first segment in the AS_PATH attribute. Enabling this command prevents a misconfigured or unauthorized peer from misdirecting traffic (spoofing the local router) by advertising a route as if it was sourced from another autonomous system.
- Step 6** (Optional.) Click the **Show Disabled** link above the object body to add all other possible configuration lines. You can enable the following options by clicking the + to the left of the option.
- **bgp asnotation dot**. Changes the default display and regular expression match format of BGP 4-byte autonomous system numbers from asplain (decimal values) to asdot (dot notation). The system uses asplain as the default display format for autonomous system numbers, but you can configure 4-byte autonomous system numbers in both the asplain and asdot format even if you do not enable this command.

In addition, the default format for matching 4-byte autonomous system numbers in regular expressions is asplain, so you must ensure that any regular expressions to match 4-byte autonomous system numbers are written in the asplain format if you do not enable this command.
 - **bgp scan time 60**. Click the number and enter the scanning interval of BGP routers for next hop validation, from 5 to 60 seconds. The default is 60 seconds.
 - **configure nexthop trigger state**. Click *state* and select either **enable** or **disable**. BGP next-hop address tracking is event driven. BGP prefixes are automatically tracked as peering sessions are established. Next-hop changes are rapidly reported to BGP as they are updated in the routing information base (RIB). This optimization improves overall BGP convergence by reducing the response time to next-hop changes for routes installed in the RIB. When a best-path calculation is run in between BGP scanner cycles, only the changes are processed and tracked. If you enable next hop address tracking, the following commands

are added. Note that if you do not configure the general options in a new object, the default is to enable this feature.

- **bgp nexthop trigger enable.** BGP next-hop address tracking improves BGP response time significantly. However, unstable Interior Gateway Protocol (IGP) peers can introduce instability to BGP. We recommend that you aggressively dampen unstable IGP peering sessions to mitigate the possible impact to BGP.
- **bgp nexthop trigger delay 5.** Click the number to change the delay interval between routing table walks for BGP next-hop address tracking. You can increase the performance of BGP next-hop address tracking by tuning the delay interval between full routing table walks to match the tuning parameters for the IGP. The default delay interval is 5 seconds, which is an optimal value for a fast-tuned IGP. In the case of an IGP that converges more slowly, you can change the delay interval to 20 seconds or more, depending on the IGP convergence time. You can set the delay from 0 to 100 seconds.
- **bgp aggregate-timer 30.** Click the number to set the interval at which BGP routes will be aggregated, from 6 to 60 seconds. The default is 30 seconds.
- **bgp router-id *router-id*.** Click *router-id* and enter the IPv4 address that should be used as the global router ID. This ID is used for any BGP process in a virtual router that does not itself specify a router ID. If you do not enable this command, the router ID is set to the highest IP address on a physical interface that is assigned to the virtual router. Use this command to ensure that the router ID remains stable.
- **bgp maxas-limit *value*.** Click *value* and enter the maximum number of autonomous system numbers in the AS-path attribute of the BGP Update message, ranging from 1 to 254. An AS-path attribute is a sequence of intermediate AS numbers between source and destination routers that form a directed route for packets to travel. The system discards routes that have a number of autonomous systems in the AS-path that exceed the specified value. In addition to setting the limit on the number of autonomous system numbers within the AS-path segment, the command limits the number of AS-path segments to ten. If you do not enable this command, no routes are discarded.

Step 7 (Optional.) Configure BGP advanced options.

Click the **Show Disabled** link if necessary to expose the following command. When editing the settings, both the **timers** and **bestpath** option sets are shown, as they have some defaults that are enabled even if you do not explicitly set them.

configure bgp advanced *advanced-option*

Click *advanced-option* and select one of the following. You can configure all of these options by clicking ... in the left column and selecting **Duplicate**.

- **timers.** Configures the timers used when communicating with BGP neighbor routers.

timers bgp 60 180 0

- First value (default 60): **Keepalive Interval.** Click the number and enter the frequency with which the system sends keepalive messages to its BGP neighbor, from 0 to 65535 seconds. We recommend that you do not specify 20 or less, or you might find routes flapping unnecessarily.
- Second value (default 180): **Hold Time.** Click the number and enter how long the system should wait after not receiving a keepalive message before declaring a BGP neighbor dead, from 0 to 65535 seconds.

- Third value (default 0): **Minimum Hold Time**. Click the number and specify the minimum acceptable hold-time configured on the BGP neighbor. The minimum acceptable hold-time must be less than, or equal to, the interval specified as the hold time for this system. The range is from 0 to 65535 seconds.
- **bestpath**. Configures options that are used in the BGP best path selection algorithm. The **bgp default local-preference** command is configured by default, but you can add the other commands by clicking the + for the command.
 - **bgp default local-preference 100**. Click the number and enter a value that indicates the preference of this system relative to other routers in the BGP AS, from 0 and 4294967295. The default value is 100. Higher values indicate higher preference. This preference is sent to all routers and access servers in the local autonomous system. This attribute is exchanged only between iBGP peers and is used to determine local policy.
 - **bgp always-compare-med**. Allow the comparison of Multi Exit Discriminator (MED) for paths from neighbors in different autonomous systems. By default, the system does not compare the MED for paths from neighbors in different autonomous systems.
 - **bgp bestpath compare-routerid**. Use the router ID as the tie breaker for best path selection when two identical routes are received from two different peers (all the attributes are the same except for the router ID). When this command is enabled, the lowest router ID will be selected as the best path when all other attributes are equal. Otherwise, the first route received is used.
 - **bgp deterministic-med**. Select the best MED path advertised from the neighboring AS.
 - **bgp bestpath med missing-as-worst**. Set a path with a missing MED attribute as the least preferred path. By default, the system considers the route with a missing MED to be the best route.
- **graceful-restart**. Configure graceful restart for the systems in a high availability or cluster configuration.
 - **bgp graceful-restart**. Enables graceful restart for non-stop forwarding. With graceful restart, the system can advertise the ability to maintain the forwarding state for an address group during restart.
 - **bgp graceful-restart restart-time 120**. Click the number and enter the maximum time period that the system will wait for a graceful-restart-capable neighbor to return to normal operation after a restart event occurs, from 1 to 3600 seconds. The default is 120 seconds.
 - **bgp graceful-restart stalepath-time 360**. Click the number and enter the maximum time period that the system will hold stale paths for a restarting peer, from 1 to 3600 seconds. All stale paths are deleted after this timer expires. The default value is 360 seconds.

Step 8 Click **OK**.

Configure the BGP Process

After you configure the BGP global settings, you can configure the BGP process. If you are using virtual routers, you can configure a separate process for each virtual router. You can configure at most one BGP process for the system or per virtual router.

Procedure

-
- Step 1** Click **Device**, then click the **Routing** summary.
- Step 2** If you enabled virtual routers, click the view icon () for the router in which you are configuring BGP.
- Step 3** Click the **BGP** tab.
- Step 4** Do one of the following:
- To create a new process, click + or click the **Create BGP Object** button.
 - Click the edit icon () for the object you want to edit. Note that when you edit an object, you might see lines that you did not directly configure. These lines are exposed to show you the default values that are being configured.

If you no longer need a process, click the trash can icon for the object to delete it.

- Step 5** Enter a name, and optionally, a description for the object.
- Step 6** Configure the minimum settings for the process:
- **router bgp** *as-number*. Click *as-number* and enter the same autonomous system (AS) number for the BGP process that you specified for the global settings. The AS number can be from 1 to 4294967295 or from 1.0 to 65535.65535. The AS number is a uniquely assigned value that identifies each network on the Internet. The system supports asplain and asdot notation as defined in RFC 5396.
 - **configure address-family** *ip-protocol*. Click *ip-protocol* and select IPv4 or IPv6. If you are using virtual routers, you can configure IPv6 for the global router only. You can configure IPv4 for any virtual router. Selecting an option adds the **address-family ipv4 unicast** or **address-family ipv6 unicast** commands, plus the following command which you must configure:
 - **configure address-family {ipv4 | ipv6} settings**. Click settings and select either **general** or **advanced**. You must configure at least one command under these options to have a minimal process, but that will not be sufficient for a meaningful process.

- Step 7** Click **Show Disabled** and customize the process to function correctly in your network.

So long as you configure a minimum set of commands, as explained above, you can save the object and customize the process settings later. The following topics explain the various sets of options. At minimum, you should configure the network settings to identify the networks for which the process will distribute routes. Both the general and advanced settings have command defaults that are appropriate in most cases.

- [Configure BGP General Settings, on page 9](#)
- [Configure BGP Advanced Settings, on page 10](#)
- [Configure the Networks for BGP to Advertise, on page 12](#)
- [Configure BGP Route Injection, on page 13](#)
- [Configure BGP Aggregate Address Settings, on page 13](#)
- [Configure BGP Filter Settings for IPv4, on page 15](#)
- [Configure BGP Neighbors, on page 16](#)

- [Configure BGP Route Redistribution from Other Routing Protocols, on page 23](#)

Step 8 (Optional.) Configure the router ID for this process.

You can configure the router ID to use for the BGP process in the BGP global settings. You can optionally configure it in the process object, too. Any router ID configured in the process object overrides the global router ID. This makes it easy to override the global value for specific virtual routers.

Click **Show Disabled** if the following command is not displayed, and click the + next to it to enable the command.

- **bgp router-id** *router-id*. Click *router-id* and enter the IPv4 address that should be used as the router ID for this process. If you do not enable this command, the router ID is set to the global router ID, or to the highest IP address on a physical interface that is assigned to the virtual router. Use this command to ensure that the router ID remains stable.

Step 9 Click **OK**.

Configure BGP General Settings

The general settings define administrative distances, timers, and for IPv4 only, next hop address tracking. These options have defaults appropriate for most networks.

Procedure

Step 1 Click **Device**, then click the **Routing** summary.

Step 2 If you enabled virtual routers, click the view icon () for the router in which you are configuring BGP.

Step 3 Click the **BGP** tab.

Step 4 Add or edit the BGP process object.

Step 5 Find the **configure address-family ipv4** or **ipv6** line. If the **general** option is already selected, move to the next step. However:

- If the *settings* variable is still displayed, click it and select **general**.
- If you have already configured advanced options, click the ... button to the left of the command and select **Duplicate**. Then, click *settings* and select **general**.

Step 6 Configure the following commands:

- **distance bgp 20 200 200**. Configures the administrative distances for BGP, from 1 to 255. These numbers are relative to the administrative values assigned to other routing processes when the system chooses the best routes. In general, the higher the value, the lower the trust rating. Use this command if another protocol is known to be able to provide a better route to a node than was actually learned via external BGP (eBGP), or if some internal routes should be preferred by BGP. Routes with a distance of 255 are not installed in the routing table. The numbers mean the following:
 - First value (default 20): **External Distance**. Click the number and enter the administrative distance for external BGP routes. Routes are external when learned from an external autonomous system.

- Second value (default 200): **Internal Distance**. Click the number and enter the administrative distance for internal BGP routes. Routes are internal when learned from peers in the local autonomous system. Changing the administrative distance of internal BGP routes is considered dangerous and is not recommended. Improper configuration can introduce routing table inconsistencies and break routing.
- Third value (default 200): **Local Distance**. Click the number and enter the administrative distance for local BGP routes. Local routes are for those networks listed with a **network** command in the BGP routing process, that is, the networks that the process is advertising, or for the networks that are being redistributed to BGP from another process.

Step 7 Click **OK**.

Configure BGP Advanced Settings

Use the advanced settings to configure a variety of options that are needed only under special circumstances. Most of these options are disabled by default.

Before you begin

If you intend to configure the **table-map** command, you must first go to the **Device > Advanced Configuration** page and create the Smart CLI route map object needed by the command.

Procedure

-
- Step 1** Click **Device**, then click the **Routing** summary.
- Step 2** If you enabled virtual routers, click the view icon () for the router in which you are configuring BGP.
- Step 3** Click the **BGP** tab.
- Step 4** Add or edit the BGP process object.
- Step 5** Find the **configure address-family ipv4** or **ipv6** line. If the **advanced** option is already selected, move to the next step. However:
- If the *settings* variable is still displayed, click it and select **advanced**.
 - If you have already configured general options, click the ... button to the left of the command and select **Duplicate**. Then, click *settings* and select **advanced**.
- Step 6** Configure the following commands. You need to click **Show Disabled** to see all but the first command when initially creating the object.
- Click + for a command to enable it.
- **bgp redistribute-internal**. Configure iBGP redistribution into an interior gateway protocol (IGP), such as EIGRP or OSPF. Exercise caution when redistributing iBGP into an IGP. Use IP prefix-list and route-map statements to limit the number of prefixes that are redistributed. Redistributing an unfiltered BGP routing table into an IGP can have a detrimental effect on normal IGP network operation. This command is enabled by default, so you need to click the - button for it to turn it off.

- **bgp suppress-inactive.** Prevent routes that are not installed in the RIB (inactive routes) from being advertised to peers. By default, BGP advertises inactive routes. Note that BGP marks routes that are not installed into the RIB with a RIB-failure flag. This flag will also appear in the output of the **show bgp** command; for example, Rib-Failure (17). This flag does not indicate an error or problem with the route or the RIB.
- **auto-summary.** (IPv4 only.) Automatically summarize subnet routes into network-level routes. Route summarization reduces the amount of routing information in the routing tables. Disable automatic summarization if you must perform routing between disconnected subnets. When automatic summarization is disabled, subnets are advertised.
- **synchronization.** Enable synchronization between BGP and your Interior Gateway Protocol (IGP) system, such as OSPF. Usually, a BGP speaker does not advertise a route to an external neighbor unless that route is local or exists in the IGP. This feature allows routers and access servers within an autonomous system to have the route before BGP makes it available to other autonomous systems. Use this command if other routers in the autonomous system do not speak BGP.
- **table-map route-map options.** (IPv4 only.) Apply a route map that sets metrics, a tag value, or a traffic index for routes that are updated in the BGP routing table, or controls whether routes are downloaded to the RIB. Click *route-map* and select the Smart CLI object that defines the route map. In the route map, you can use match clauses for IP access list, autonomous system paths, communities, prefix lists, and next hop.

You can determine how the route map is used by clicking *options* and choosing either blank or **filter**:

- If you do not select **filter**, the system uses the route map to set certain properties of a route before the route is installed into the RIB. The route is always downloaded, regardless of whether it is permitted or denied by the route map.
- If you select **filter**, the route map also controls whether the BGP route is downloaded to the RIB. Only routes permitted in the route map are downloaded; denied routes are not downloaded.
- **default-information originate.** Configure BGP to advertise a default route (network 0.0.0.0). The configuration of the **default-information originate** command is similar to the configuration of the **network** command. The **default-information originate** command, however, requires explicit redistribution of the route 0.0.0.0, which you must also configure in this object. The **network** command requires only that the route 0.0.0.0 be present in the Interior Gateway Protocol (IGP), such as OSPF, routing table. For this reason, the **network** command is preferred for distributing a default route.
- **maximum paths 1.** Control the maximum number of parallel BGP routes that can be installed in a routing table, from 1 to 8. Use this command to configure equal-cost or unequal-cost multipath load sharing for BGP peering sessions. In order for a route to be installed as a multipath in the BGP routing table, the route cannot have a next hop that is the same as another route that is already installed. The BGP routing process will still advertise a best path to BGP peers when BGP multipath load sharing is configured. For equal-cost routes, the path from the neighbor with the lowest router ID is advertised as the best path.

To configure BGP equal-cost multipath load sharing, all path attributes must be the same. The path attributes include weight, local preference, autonomous system path (the entire attribute and not just the length), origin code, Multi Exit Discriminator (MED), and Interior Gateway Protocol (IGP) distance.
- **maximum paths ibgp 1.** Control the maximum number of internal BGP routes that can be installed to the routing table, from 1 to 8. The considerations for multipath iBGP are the same as described under the **maximum paths** command above.

Step 7 Click **OK**.

Configure the Networks for BGP to Advertise

You need to define the networks to be advertised by the BGP routing process.

Before you begin

Create the network objects that define the networks to advertise. You can define IPv4 or IPv6 networks, or both, depending on the address families you configure for BGP.

If the network objects specify large network spaces, you can also create route maps to apply against the network object to filter out subnets within the larger space that you do not want to advertise. Only routes that match the route map specifications are advertised. Use Smart CLI to create the route map object.

Procedure

Step 1 Click **Device**, then click the **Routing** summary.

Step 2 If you enabled virtual routers, click the view icon () for the router in which you are configuring BGP.

Step 3 Click the **BGP** tab.

Step 4 Add or edit the BGP process object.

The network commands are within the command sets beneath the **configure address family ipv4** or **ipv6** command. You must configure the address family to configure the networks to advertise.

The **network** command within each address group must specify addresses that match the address family you are configuring.

Step 5 Click **Show Disabled** to expose all commands, then click + to enable the **network** or **network route-map** command, and configure the options:

- **network-object**. Click the variable and select the network object that defines the networks to advertise: IPv4 network address and mask or IPv6 network address and prefix.
- **route-map map-tag**. Click the variable and select the route map that should be applied to the network object to filter which addresses within the range should be advertised.
- (Optional; IPv6 only.) **prefix-name**. Click the variable and enter the name of a DHCPv6 prefix to advertise the prefix. If you configure this option, the network object acts as a subnet for the prefix. To use this option, you must enable the DHCPv6 Prefix Delegation client, which requires that you use FlexConfig to add the **ipv6 dhcp client pd** command to an interface in interface configuration mode.

Step 6 You can click ... > **Duplicate** next to the **network** or **network route-map** command to configure additional networks to advertise.

Step 7 Click **OK**.

Configure BGP Route Injection

You can configure conditional route injection to inject more specific routes into a BGP routing table. Conditional route injection allows you to originate a more specific prefix into a BGP routing table without a corresponding match. A valid parent route must exist for any injected prefixes. Only prefixes that are equal to or more specific than the aggregate route (existing prefix) can be injected.

Before you begin

You must create the route maps needed to define the prefixes. These route maps must meet the requirements explained in the procedure.

Procedure

-
- Step 1** Click **Device**, then click the **Routing** summary.
- Step 2** If you enabled virtual routers, click the view icon () for the router in which you are configuring BGP.
- Step 3** Click the **BGP** tab.
- Step 4** Add or edit the BGP process object.
- The route injection commands are within the command sets beneath the **configure address family ipv4** or **ipv6** command. You must configure the address family to configure the networks to advertise.
- Step 5** Click **Show Disabled** to expose all commands, then click + to enable the **bgp inject-map** command.
- Step 6** Configure the command properties:
- **inject-map** *inject-map*. Click the variable and select the route map that defines the prefixes that will be created and installed into the routing table. Injected prefixes are installed in the local BGP RIB. A valid parent route must exist; Only prefixes that are equal to or more specific than the aggregate route (existing prefix) can be injected. The route map must use a prefix list to specify the routes to be injected.
 - **exist-map** *exist-map*. Click the variable and select the route map that defines the prefix that the BGP speaker will track. This route map must use prefix lists to specify the aggregate prefix and also the route source. The route source would be a router, for example, 10.2.1.1/32, rather than a subnet.
 - **options**. Optionally, click the variable and select **copy-attributes**. This option configures the injected prefix to inherit the same attributes as the aggregate route. If you do not select this keyword, the injected prefix will use the default attributes for locally originated routes.
- Step 7** You can click ... > **Duplicate** next to the **bgp inject-map** command to configure additional route injection rules.
- Step 8** Click **OK**.
-

Configure BGP Aggregate Address Settings

BGP neighbors store and exchange routing information and the amount of routing information increases as more BGP speakers are configured. Route aggregation is the process of combining the attributes of several different routes so that only a single route is advertised. Aggregate prefixes use the classless interdomain routing (CIDR) principle to combine contiguous networks into one classless set of IP addresses that can be summarized in routing tables. As a result fewer routes need to be advertised.

If you configure an aggregate route with no keywords on the command, the system will create an aggregate entry in the BGP routing table if any more-specific BGP routes are available that fall within the specified range. (A longer prefix that matches the aggregate must exist in the Routing Information Base (RIB).) The aggregate route will be advertised as coming from your autonomous system and will have the atomic aggregate attribute set to show that information might be missing. The atomic aggregate attribute is set unless you specify the **as-set** keyword.

The following procedure explains how to configure the aggregation of specific routes into one route.

Before you begin

If you want to apply route maps to fine-tune which routes are aggregated or the attributes set on the aggregate route, create the Smart CLI route map objects.

Procedure

-
- Step 1** Click **Device**, then click the **Routing** summary.
- Step 2** If you enabled virtual routers, click the view icon () for the router in which you are configuring BGP.
- Step 3** Click the **BGP** tab.
- Step 4** Add or edit the BGP process object.
- The aggregation commands are within the command sets beneath the **configure address family ipv4** or **ipv6** command. You must configure the address family to configure aggregation.
- Step 5** Click **Show Disabled** to expose all commands, then click + to enable the **configure aggregate-address** command.
- Step 6** Click the *map-type* variable and select which types of route map you want to apply to this particular aggregate route.
- This option simply determines which parameters are included on the **aggregate-address** command that will be added to the object. You can apply up to 3 separate route maps, to suppress routes from the aggregation, to advertise routes, and to define attributes to apply to the aggregate route.
- Select **no-map** if you do not need to apply any route maps.
 - Select **all** if you want to apply route maps for all three options.
 - Select the appropriate keyword combinations if you want to apply one or two maps, but not all: **suppress-map**, **advertise-map**, **attribute-map**, **suppress-advertise**, **suppress-attribute**, **advertise-attribute**.
- Step 7** Configure the properties of the route to be aggregated.
- Following is a complete list of the properties. What you see depends on your map type selection.
- **network-object**. Click the variable and select the network object that defines the address space that you want to aggregate. The object must use IPv4 or IPv6 addressing that matches the address type you are configuring. For example, you could aggregate routes for all 10.0.0.0/8 subnets.
 - **suppress-map** *suppress-route-map*. Click the variable and select the route map to suppresses the advertisement of specified routes. You can use the match clauses of the route map to selectively suppress

some more-specific routes of the aggregate and leave others unsuppressed. The route map can match routes based on access lists and autonomous system paths.

- **advertise-map** *advertise-route-map*. Click the variable and select the route map that selects specific routes that will be used to build different components of the aggregate route, such as AS_SET or community. This is useful when the components of an aggregate are in separate autonomous systems and you want to create an aggregate with AS_SET, and advertise it back to some of the same autonomous systems. You must remember to omit the specific autonomous system numbers from the AS_SET to prevent the aggregate from being dropped by the BGP loop detection mechanism at the receiving router. The route map can match routes based on access lists and autonomous system paths.
- **attribute-map** *attribute-route-map*. Click the variable and select the route map that changes attributes of the aggregate route. This is useful when one of the routes forming the AS_SET is configured with an attribute such as the community no-export attribute, which would prevent the aggregate route from being exported.
- **options**. Click the variable and select one, all, or none of the following options:
 - **as-set**. Generate autonomous system set path information for the aggregate route. The path advertised for this route will be an AS_SET consisting of all elements contained in all paths that are being summarized. Do not use this keyword when aggregating many paths, because this route must be continually withdrawn and updated as autonomous system path reachability information for the summarized routes changes.
 - **summary-only**. Suppress advertisements of more-specific routes to all neighbors.

- Step 8** You can click ... > **Duplicate** next to the **configure aggregate-address** command to configure additional routes to aggregate.
- Step 9** Click **OK**.

Configure BGP Filter Settings for IPv4

You can create filter rules to restrict the routing information the system learns from, or advertises to, other routing protocols.

The configuration explained here applies to all local processes and for filtering updates to all BGP neighbors. You can configure different filtering rules per neighbor in the neighbor settings.

Before you begin

Create the Smart CLI standard access list objects you need for each filter rule. Use deny access control entries (ACEs) to filter out routes that match the entry, and permit ACEs for the routes that should be updated.

Procedure

- Step 1** Click **Device**, then click the **Routing** summary.
- Step 2** If you enabled virtual routers, click the view icon () for the router in which you are configuring BGP.
- Step 3** Click the **BGP** tab.
- Step 4** Add or edit the BGP process object.

The filtering commands are within the command sets beneath the **configure address family ipv4** command. You must configure the address family to configure filtering. These rules are not available for IPv6.

- Step 5** Click **Show Disabled** to expose all commands, then click + to enable the **configure filter-rules direction** command.
- Step 6** Click *direction* and select **in**, for filtering incoming updates, or **out**, for filtering outbound updates.
- Step 7** For inbound filters, you can optionally specify the interface on which to filter updates. If you do not specify an interface, the filter applies to all updates received on any interface.
- Click + to enable the **distribute-list acl-name in interface interface** command.
 - Click the *interface* variable and select the interface.
- Step 8** For outbound filters, you can optionally specify the protocol, to limit the filter to routes advertised to that routing process.
- There are two forms of the **distribute-list out** command, one with an *identifier* variable after the *protocol* variable, and one without the identifier. You can select the following protocols, but they are divided between these command versions based on whether you must provide the additional identifier information.
- **connected**. For routes established for networks that are directly connected to the system's interfaces.
 - **static**. For static routes you manually created.
 - **rip**. For routes advertised to RIP.
 - **bgp autonomous-system**. For routes advertised to BGP. Click *identifier* and enter the autonomous system number for the BGP process defined on the system.
 - **eigrp autonomous-system**. For routes advertised to EIGRP. Click *identifier* and enter the autonomous system number for the EIGRP process defined on the system.
 - **ospf process-id**. For routes advertised to OSPF. Click *identifier* and enter the process ID for the OSPF process defined on the system.
- Step 9** You can click ... > **Duplicate** next to the **configure filter-rules** command to define another filter rule. Define as many as you require.
- Step 10** Click **OK**.

Configure BGP Neighbors

You need to define the neighbors with whom BGP will exchange routing updates.

Before you begin

Several optional commands require Smart CLI objects, for route maps, prefix lists, and so forth. Examine the options you need to configure to determine if you need objects. You must create the Smart CLI objects before configuring the associated BGP command.

Procedure

- Step 1** Click **Device**, then click the **Routing** summary.

- Step 2** If you enabled virtual routers, click the view icon () for the router in which you are configuring BGP.
- Step 3** Click the **BGP** tab.
- Step 4** Add or edit the BGP process object.
- The neighbor commands are within the command sets beneath the **configure address family ipv4** or **ipv6** command. You must configure the neighbors separately for each address family.
- Step 5** Click **Show Disabled** to expose all commands, then click + to enable the **configure neighbor** command.
- Step 6** Configure the basic neighbor parameters on the neighbor command:
- **neighbor** *neighbor-address*. Click the variable and enter the IPv4 or IPv6 address of the BGP neighbor router, as appropriate for the address group you are configuring.
 - **remote-as** *as-number*. Click the variable and enter the autonomous system number of the BGP neighbor router.
 - **config-options**. Click the variable and select **properties**. The only property that is configured by default activates the neighbor. You can adjust other options as explained in this procedure.
- Step 7** (Optional.) Configure the neighbor general settings.
- a) Click + to enable the **configure neighbor** *neighbor-address* **remote-as** *settings* command. Click **Show Disabled** if you cannot see the command.
 - b) Click *settings* and select **general**.
 - c) On the **configure neighbor description** command, either click the variable and enter a description of the neighbor (such as its location or purpose), up to 80 characters, or click - to disable the command if you do not want a description. A description cannot include spaces or question marks.
 - d) (IPv4 only.) The **configure neighbor shutdown** command is initially enabled. This command disables communication with this BGP neighbor, terminating any active session and removing all associated routing information. If you want to actively communicate with this neighbor, click - to disable this command.
 - e) On the **configure neighbor fall-over bfd** command, either click *option* and select **single-hop** or **multi-hop** (based on your BFD configuration), or click - to disable the command.
- This command registers BGP to receive forwarding path detection failure messages from Bidirectional Forwarding Detection (BFD). Whether you select single-hop or multi-hop depends on the type of BFD template you created and attached to the interface that faces this neighbor. Ensure that your selection here is consistent with the BFD template. You must use FlexConfig to build and apply BFD templates.
- Step 8** (Optional.) Configure the neighbor advanced settings.
- a) If you have already configured it, click ... > **Duplicate** for the **configure neighbor** *neighbor-address* **remote-as** *settings* command, or simply click + to enable it if it is not yet in use. Click **Show Disabled** if you cannot see the command.
 - b) Click *settings* and select **advanced**.
 - c) On the **neighbor password** command, either click the *secret* variable and select the secret key object that contains the password to use when authenticating the neighbor, or click - to disable the command if you do not want use message digest 5 (MD5) authentication. You can create the key object while editing the BGP object.
- The secret key object must contain a key case-sensitive password of up to 25 characters in length. The string can contain any alphanumeric characters, including spaces, and the special characters ` ~ ! @ # \$ % ^ & * () - _ = + | \ }] { [" ' : ; / > < , . ? . However, you cannot specify a password in the format number-space-anything. The space after the number can cause authentication to fail.

Ensure the neighbor is configured to use the same password.

- d) On the **configure neighbor hops** command, click the *options* variable and select one of the following, or click - to disable the command if the peer is not multiple hops away (that is, not directly connected to this system). Use these options with care as you can end up with routing loops and oscillating routes: configuring only directly-connected peers is preferred.

- **ebgp-multihop**. Accept and attempt BGP connections to external peers residing on networks that are not directly connected. If you select this option, the following commands are added:
 - **neighbor ebgp-multihop 255**. Click the 255 and enter the time-to-live value in number of hops, from 1-255.
 - **neighbor disable-connected-check**. Click + to enable this command to disable connection verification to establish an eBGP peering session with a single-hop peer that uses a loopback interface. Without this command, if the peer is not directly connected to the same network segment, connection verification will prevent the peering session from being established.

- **ttl-security-hop**. Secure a BGP peering session and configure the maximum number of hops that separate two external BGP (eBGP) peers. If you select this option, the following command is added:

neighbor ttl-security hops hop-count. Click the variable and enter the maximum number of hops that separate the peers, from 1-254.

The **neighbor ttl-security** command provides a lightweight security mechanism to protect BGP peering sessions from CPU utilization-based attacks. These types of attacks are typically brute force Denial of Service (DoS) attacks that attempt to disable the network by flooding the network with IP packets that contain forged source and destination IP addresses in the packet headers.

This feature leverages designed behavior of IP packets by accepting only IP packets with a TTL count that is equal to or greater than the locally configured value. Accurately forging the TTL count in an IP packet is generally considered to be impossible. Accurately forging a packet to match the TTL count from a trusted peer is not possible without internal access to the source or destination network.

To maximize the effectiveness of this feature, the hop-count value should be strictly configured to match the number of hops between the local and external network. However, you should also take path variation into account when configuring this feature for a multihop peering session. Ensure that you configure this feature on all routers in the network.

- e) On the **neighbor version** command, click the *version-number* variable and enter 4 to force the software to use BGP version 4, or click - to disable the command. The software uses version 4 by default and negotiates down to version 2 dynamically if necessary. Configuring 4 on this command prevents version negotiation.
- f) On the **neighbor transport connection-mode** command, click the *options* variable and select whether the TCP connection is **active** or **passive**, or click - to disable the command and leave the mode to default.
- g) On the **neighbor transport path-mtu-discovery** command, click the *options* variable and select **blank** to enable path MTU discovery, or **disable** to disable it. Selecting blank is the same as clicking - to disable the command, as the system performs path MTU discovery by default. Path MTU discovery makes it possible for the BGP session to take advantage of larger MTU links.

Step 9

(Optional.) Configure the neighbor migration settings.

The migration settings configure the **neighbor local-as** command. The **neighbor local-as** command is used to customize the AS_PATH attribute by adding and removing autonomous system numbers for routes received

from eBGP neighbors. The configuration of this command allows a router to appear to external peers as a member of another autonomous system for the purpose of autonomous system number migration. This feature simplifies the process of changing the autonomous system number in a BGP network by allowing the network operator to migrate customers to new configurations during normal service windows without disrupting existing peering arrangements.

You can perform this migration for only true eBGP peering sessions. This command does not work for two peers in different sub-autonomous systems of a confederation.

Caution

BGP prepends the autonomous system number from each BGP network that a route traverses to maintain network reachability information and to prevent routing loops. Configure this command only for autonomous system migration, and de-configure it after the transition has been completed. This procedure should be attempted only by an experienced network operator. You can create routing loops through improper configuration.

- a) If you have already configured it, click ... > **Duplicate** for the **configure neighbor** *neighbor-address remote-as settings* command, or simply click + to enable it if it is not yet in use. Click **Show Disabled** if you cannot see the command.
- b) Click *settings* and select **migration**. This adds the following command:
configure neighbor-address local-as local-as-number options
- c) Click the *local-as-number* variable and enter the local autonomous system (AS) number to prepend to the AS_PATH attribute, from 1 to 4294967295 (asplain notation) or from 1.0 to 65535.65535 (asdot notation). You cannot specify the autonomous system number from the local BGP routing process or from the network of the remote peer.
- d) Click the *options* variable and select one of the following. Note that selecting an item in this list (other than **none**) also selects all options above it in the list. This is expected: the options are not truly independent.
 - **none**. Do not configure any of the following options.
 - **no-prepend**. Do not prepend the local autonomous system number to any routes received from the eBGP neighbor.
 - **replace-as**. Replace the real autonomous system number with the local autonomous system number in the eBGP updates. The autonomous system number from the local BGP routing process is not prepended.
 - **dual-as**. Configure the eBGP neighbor to establish a peering session using the real autonomous system number (from the local BGP routing process) or by using the local autonomous system number.

Step 10 (Optional. IPv4 only.) Configure the neighbor high availability (HA) settings.

The HA mode settings configure the **neighbor ha-mode graceful-restart** command, which enables or disables the graceful restart capability for an individual BGP neighbor. Use the disable keyword to disable the graceful restart capability when graceful restart has been previously enabled for the BGP peer.

The graceful restart capability is negotiated between nonstop forwarding (NSF)-capable and NSF-aware peers in OPEN messages during session establishment. If you enable the graceful restart capability after a BGP session has been established, you will need to restart the session with a soft or hard reset.

The HA mode setting configures graceful restart for an individual neighbor. Instead, you can use the BGP global settings to enable graceful restart for all neighbors.

- a) If you have already configured it, click ... > **Duplicate** for the **configure neighbor neighbor-address remote-as settings** command, or simply click + to enable it if it is not yet in use. Click **Show Disabled** if you cannot see the command.
- b) Click *settings* and select **ha-mode**.
- c) If you want to disable graceful restart, click *options* on the **neighbor ha-mode graceful-restart** command and select **disable**. Select blank to reverse a previous disable action.

Step 11 (Optional.) Configure neighbor activation options.

When you configure a new neighbor, it is activated by default. You need to enable the activation settings if you want the neighbor to be disabled initially, or to configure other activation settings.

- a) Click + to enable the **configure neighbor neighbor-address activate activate-options** command. Click **Show Disabled** if you cannot see the command.
- b) Click *activate-options* and select **properties**.
- c) The **neighbor neighbor-address activate** command is added in the enabled state. Click - to disable the command and configure the neighbor as initially disabled. You will need to edit this object to enable the neighbor when you are ready to communicate with it.

Step 12 (Optional.) Configure filtering in the neighbor activation settings.

- a) If you have already configured it, click ... > **Duplicate** for the **configure neighbor neighbor-address activate settings** command, or simply click + to enable it if it is not yet in use. Click **Show Disabled** if you cannot see the command.
- b) Click *settings* and select **filtering**.
- c) Configure filtering to control the prefixes received from or sent to this neighbor using any combination of the following neighbor commands. Click - to disable any that you do not want to use. All of these commands allow filtering in both the inbound and outbound direction: click ... > **Duplicate** for a command if you want to configure both directions.

Do not apply both a **neighbor distribute-list** and a **neighbor prefix-list** command to a neighbor in the same direction. These two commands are mutually exclusive, and only one of them can be applied to each inbound or outbound direction.

- **distribute-list** *acl options*. (IPv4 only.) Filter prefixes based on the selected standard access list (ACL). Then, click *options* and select whether to apply the filter in the **in** or **out** direction.
 - **route-map** *route-map options*. Filter prefixes based on the selected route map. Then, click *options* and select whether to apply the filter in the **in** or **out** direction. Within the route map, you can configure filtering based on access list, AS path, prefix, and distribution lists.
 - **prefix-list** *prefix-list options*. Filter prefixes based on the selected IPv4 or IPv6 prefix list. Then, click *options* and select whether to apply the filter in the **in** or **out** direction.
 - **filter-list** *as-path options*. Filter prefixes based on the selected AS Path filter object. Then, click *options* and select whether to apply the filter in the **in** or **out** direction.
- d) On the **configure prefix-limit neighbor neighbor-address limit-options** command, either click *limit-options* and select one of the following, or click - to disable the command. Selecting any option adds some form of the **neighbor maximum-prefix** command, with additional options you need to configure. Use this command to control how many prefixes can be received from the neighbor.
 - **none**. Configure the basic form of the command without additional parameters. Click the variables and configure the following values:

- **max-prefix-limit.** The maximum number of prefixes allowed from this neighbor, from 1 - 2147483647. You must also configure this variable if you select any of the other options.
- **75 (threshold).** The percentage of the maximum at which the router starts to generate a warning message, from 1 to 100. The default is 75%.
- **restart.** Stop the peering session with the neighbor when the limit is reached. Click the *restart-interval* variable and configure how long the system should wait before restarting the session, from 1 to 65535 minutes.
- **warning-only.** Do not stop the session when the limit is reached. Instead, simply issue a warning syslog message and continue the session.

Step 13 (Optional.) Configure routes in the neighbor activation settings.

- If you have already configured it, click ... > **Duplicate** for the **configure neighbor** *neighbor-address* **activate** *settings* command, or simply click + to enable it if it is not yet in use. Click **Show Disabled** if you cannot see the command.
- Click *settings* and select **routes**.
- On the **neighbor advertisement-interval** command, click the *value* variable and enter the minimum route advertisement interval between sending route updates to this neighbor, from 0 to 600 seconds, or click - to disable the command and leave the interval to default to 0, for iBGP and eBGP sessions in a virtual router, or 30 for eBGP sessions not in a virtual router. The value 0 means the system will send updates whenever the routing table changes, regardless of how frequent that might be.
- On the **neighbor advertise-map** command, either configure the following options to conditionally advertise selected routes to the neighbor, or click - to disable the command, thus sending all route updates to the neighbor unconditionally.

The routes (prefixes) that will be conditionally advertised are defined in two route maps: an advertise map and either an exist map or non-exist map.

The route map associated with the exist map or non-exist map specifies the prefix that the BGP speaker will track.

The route map associated with the advertise map specifies the prefix that will be advertised to the specified neighbor when the condition is met.

If you configure an exist map, the condition is met when the prefix exists in both the advertise map and the exist map.

If you configure a non-exist map, the condition is met when the prefix exists in the advertise map, but does not exist in the non-exist map.

If the condition is not met, the route is withdrawn and conditional advertisement does not occur. All routes that may be dynamically advertised or not advertised need to exist in the BGP routing table for conditional advertisement to occur.

- **advertise-route-map.** Click this variable and select the route map that defines which routes should be advertised if the conditions of the exist map or non-exist map are met.
- **options condition-route-map.** Click options and select one of the following:
 - **exist-map.** Click the variable and select the exist route map.
 - **non-exist-map.** Click the variable and select the non-exist route map.

- e) The **neighbor** *neighbor-address* **remove-private-as** command is added in the enabled state. Click - to disable the command. This command removes private autonomous system numbers from the eBGP outbound routing updates. The private AS values are 64512 to 65535.
- f) On the **configure neighbor default-originate** command, either click *options* and select one of the following, or click - to disable the command.
 - **none**. Allows the system to send a default route to the neighbor unconditionally.
 - **route-map**. Have the system conditionally send a default route to the neighbor. When used with a route map, the default route is injected if the route map contains a match IP address clause and there is a route that matches the IP access list exactly. You can use standard or extended access lists in the route map to define the default routes. You must click the *route-map* variable in **neighbor default-originate** command that is added to the object and select the route map.

Step 14 (Optional.) Configure timers in the neighbor activation settings.

If you configure timers for a neighbor, the settings override the timers configured for all BGP neighbors in the global BGP settings.

- a) If you have already configured it, click ... > **Duplicate** for the **configure neighbor** *neighbor-address* **activate** *settings* command, or simply click + to enable it if it is not yet in use. Click **Show Disabled** if you cannot see the command.
- b) Click *settings* and select **timers**.
- c) On the **neighbors timers** command, configure the following variables:
 - **keepalive-interval**. The frequency with which the system sends keepalive messages to this neighbor, from 0 to 65535 seconds. The default is 60 seconds if you do not configure this command.
 - **hold-time**. The interval after not receiving a keepalive message that the system declares this neighbor dead, 0 to 65535 seconds. The default is 180 seconds if you do not configure this command.
 - **0** (minimum hold time). The minimum acceptable hold-time that can be configured on this neighbor, from 0 to 65535 seconds. This value must be less than or equal to the hold time configured for this system. If the neighbor's hold time is less than this value, the system will not establish a BGP session with the neighbor.

Step 15 (Optional.) Configure advanced neighbor activation settings.

- a) If you have already configured it, click ... > **Duplicate** for the **configure neighbor** *neighbor-address* **activate** *settings* command, or simply click + to enable it if it is not yet in use. Click **Show Disabled** if you cannot see the command.
- b) Click *settings* and select **advanced**.
- c) Decide which of the following **neighbor** commands to leave enabled. Click - to disable unwanted options.
 - **send-community**. Sends the communities attribute to the neighbor.
 - **weight** *value*. Click the variable to assign the initial weight to routes learned from this neighbor, from 0 and 65535. If you do not configure this command, routes learned through another BGP peer have a default weight of 0 and routes sourced by the local router have a default weight of 32768. However, any route weight that is set using a route map overrides the weight configured using this command.
 - **next-hop-self**. Configures the router as the next hop for a BGP-speaking neighbor. This command is useful in unmeshed networks (such as Frame Relay or X.25) where BGP neighbors may not have direct access to all other neighbors on the same IP subnet.

- Step 16** You can click ... > **Duplicate** next to the **configure neighbor** command to define another neighbor. Define as many as you require.
- Step 17** Click **OK**.

Configure BGP Route Redistribution from Other Routing Protocols

You can control the redistribution of routes into a BGP process from other routing protocols, connected routes, and static routes.

Before you begin

It is best practice to configure the routing process from which you will redistribute routes, and deploy your changes, before you configure redistribution into BGP.

If you want to apply a route map to fine-tune which routes are redistributed, create the Smart CLI route map object. Routes that match the route map are redistributed, and all non-matching routes are not redistributed.

Procedure

- Step 1** Click **Device**, then click the **Routing** summary.
- Step 2** If you enabled virtual routers, click the view icon () for the router in which you are configuring BGP.
- Step 3** Click the **BGP** tab.
- Step 4** Add or edit the BGP process object.
- The redistribution commands are within the command sets beneath the **configure address family ipv4** or **ipv6** command. You must configure the address family to configure redistribution.
- Step 5** Click **Show Disabled** to expose all commands, then click + to enable the **configure ipv4/ipv6 redistribution** command.
- Step 6** Click the *protocol* variable and select the source process from which you are redistributing routes. You can redistribute **connected** and **static** routes, or routes generated by **eigrp** (IPv4 only), **isis**, **ospf**, or **rip** (IPv4 only).
- Step 7** If you select a routing process, click the *identifier* variable and enter the required value:
- **eigrp**. Enter the autonomous system number.
 - **ospf**. Enter the process ID number.
 - **connected**, **static**, **isis**, **rip**. Enter **none**. Even if you enter a different value, it will be ignored.
- Step 8** (Optional; IS-IS only.) On the **redistribute isis level-2** command, click **level-2** and select whether you are redistributing routes learned only within an IS-IS area (**level-1**), between IS-IS areas (**level-2**) or both (**level-1-2**).
- Step 9** (Optional; all protocols.) To fine-tune the metrics for redistributed routes, click + to enable the following command and configure the options:
- redistribute protocol metric metric-value**
- Click the variable and enter the metric value for the routes being distributed, from 0 to 4294967295.

- Step 10** (Optional; all protocols.) To fine-tune which routes are redistributed based on a route map, click + to enable the **redistribute route-map** command, click the variable, and select the route map that defines your restrictions.
- If you do not apply a route map, all routes for the process (that fit the other commands configured for redistribution), are redistributed.
- Step 11** (Optional; OSPF only.) The following commands are enabled by default when you redistribute routes from an OSPF process. You can click - to disable unwanted commands.
- These commands specify the criteria by which OSPF routes are redistributed into other routing domains.
- **redistribute ospf match external 1.** Routes that are external to the autonomous system, but are imported into OSPF as Type 1 external routes.
 - **redistribute ospf match external 2.** Routes that are external to the autonomous system, but are imported into OSPF as Type 2 external routes.
 - **redistribute ospf match internal.** Routes that are internal to a specific autonomous system.
 - **redistribute ospf match nssa-external 1.** Routes that are external to the autonomous system, but are imported into OSPF as Type 1 external routes and marked as Not-So-Stubby-Area (NSSA) only.
 - **redistribute ospf match nssa-external 2.** Routes that are external to the autonomous system, but are imported into OSPF as Type 2 external routes and marked as Not-So-Stubby-Area (NSSA) only.
- Step 12** You can click ... > **Duplicate** next to the **configure redistribution** command to configure redistribution for another protocol. Configure redistribution for each protocol that makes sense for your network.
- Step 13** Click **OK**.
-

Monitoring BGP

To monitor and troubleshoot BGP, open the CLI console or log into the device CLI and use the following commands. You can also select some of these commands from the **Commands** menu on the Routing page.

Use **show bgp ?** to get lists of additional options. For example, you can specify autonomous system number, and virtual router, to limit the information you see, as well as other options to target just the information you are looking for. The following list is a summary only.

- **show bgp**
Displays the entries in the BGP routing table.
- **show bgp cidr-only**
Displays routes with non-natural network masks (that is, classless interdomain routing, or CIDR).
- **show bgp community**
Displays routes that belong to specified BGP communities.
- **show bgp community-list**
Displays routes that are permitted by the BGP community list.
- **show bgp filter-list *access-list-number***

Displays routes that conform to a specified filter list.

- **show bgp injected-paths**

Displays all the injected paths in the BGP routing table.

- **show bgp ipv4 unicast**

Displays entries in the IP version 4 (IPv4) BGP routing table for unicast sessions.

- **show bgp ipv6 unicast**

Displays entries in the IPv6 BGP routing table.

- **show bgp neighbors**

Displays information about BGP and TCP connections to neighbors.

- **show bgp paths**

Displays all the BGP paths in the database.

- **show bgp prefix-list**

Displays information about a prefix list or prefix list entries.

- **show bgp regexp *regexp***

Displays routes that match the autonomous system path regular expression.

- **show bgp rib-failure**

Displays BGP routes that failed to install in the Routing Information Base (RIB) table.

- **show bgp summary**

Display the status of all BGP connections.

- **show bgp update-group**

Display information about the BGP update groups.

