



Working with Connectivity in Your
Nexus Dashboard LAN Fabrics,
Release 4.3.1

Table of Contents

New and changed information	1
Navigate to the Connectivity page	2
Interfaces	3
Add interfaces	9
Breakout	11
UnBreakout	11
Edit interfaces	11
Additional information	14
Enable VLAN mapping	15
Limitations	16
Edit interfaces associated with links	16
Delete interfaces	17
Shut down and bring up interfaces	17
View interface configuration	18
Rediscover interfaces	18
View interface history	18
Deploy interface configurations	18
Create external fabric interfaces	19
Sync up out-of-band switch interface configurations	19
Guidelines	20
Sync up switch interface configurations	21
Interface groups	24
Create an interface group	25
Remove interfaces from an interface group	26
Attach networks to an interface group	27
Detach a network from an interface group	28
Delete an interface group	28
Links	29
Links	29
Protocol view	30
Create intra-fabric links	30
Create inter-fabric links	32
Establishing inter-fabric connectivity using VRF Lite	35
About connecting two NX-OS fabrics with MACsec using QKD	56
Routing policies	64
Navigate to the Routing policies page	64
L3Outs	64
Guidelines and limitations: L3Outs	64
Working with L3Outs	65
Create an L3Out	66
Edit an L3Out	76

Attach an L3Out	76
Detach an L3Out	76
Delete an L3Out	77
Configure a track	77
Route-maps	77
Create a route-map	78
Edit a route-map	83
Delete a route-map	83
ACLs	83
Create an ACL	84
Edit an ACL	90
Delete an ACL	91
Prefix-list	91
Create a prefix-list	91
Edit a prefix-list	92
Delete a prefix-list	93
About community lists and extended community lists	93
Community-list	94
Extended community-list	98
Inter-Fabric	103
L3 Neighbors	104
View L3 neighbors in a fabric	104
Guidelines and limitations for L3 neighbors	105
Endpoints	106
View endpoints in a fabric	106
Guidelines and limitations for endpoints	107
Routes	108
View routes in a fabric	108
View IPv4 or IPv6 routes	108
Guidelines and limitations for routes	108
Multicast traffic monitoring overview	109
Multicast routes	109
Multicast telemetry flow path visibility	109
Prerequisites	109
Monitor multicast routes	110
Guidelines and limitations for multicast routes	112
Flows	114
Flows hardware requirements	114
Flows guidelines and limitations for NX-OS fabrics	114
View flows	116
View unicast flows	117
View multicast flows	121
Virtual Infrastructure	123

View virtual machine VMs	124
View Kubernetes pods	125
Copyright	127

New and changed information

The following table provides an overview of the significant changes up to this current release. The table does not provide an exhaustive list of all changes or of the new features up to this release.

Release Version	Feature	Description
Nexus Dashboard 4.3.1	Multicast telemetry flow path visibility	Beginning with Nexus Dashboard 4.3.1, you can monitor multicast traffic with end-to-end flow path visibility from the first-hop router (FHR) to the last-hop router (LHR). Nexus Dashboard exports multicast routes and state information from Cisco N9000 Series switches, enabling you to visualize the complete forwarding path, access granular source and receiver statistics, and identify top multicast groups by traffic and receivers. For more information, see Multicast traffic monitoring overview.
Nexus Dashboard 4.3.1	Extended loopback ID range for Cisco Catalyst 9000 Series switches	Beginning with Nexus Dashboard 4.3.1, Cisco Catalyst 9000 Series switches support loopback IDs in the range 0-2147483646 for user-defined loopback interfaces. For more information, see Add interfaces.
Nexus Dashboard 4.3.1	Catalyst Center telemetry configuration support	Beginning with Nexus Dashboard 4.3.1, you can configure flow monitor entries and device tracking policy settings on Cisco Catalyst switch interfaces in Campus VXLAN EVPN fabrics and on IOS XE switch interfaces in External fabrics. These settings allow the interface intent in Nexus Dashboard to match telemetry-related configuration pushed by Catalyst Center. For more information, see Catalyst Center Assurance interface fields.

Navigate to the Connectivity page

To navigate to the **Connectivity** page:

1. Navigate to the main **Fabrics** page:

Manage > Fabrics

2. In the table showing all of the Nexus Dashboard fabrics that you have already created, locate the LAN or IPFM fabric where you want to configure connectivity.
3. Single-click on that fabric.

The **Overview** page for that fabric appears.

4. Click the **Connectivity** tab.

These subtabs provide more focused connectivity options:

- [Interfaces](#)
- [Links](#)
- [L3 Neighbors](#)
- [Endpoints](#)
- [Routes](#)
- [Inter-Fabric](#)
- [Flows](#)
- [Virtual Infrastructure](#)

Interfaces

The **Interfaces** option displays these details about the switch and its vPC peer:

- interfaces that are discovered
- intended interfaces not yet deployed, and
- Virtual Port Channels (vPCs)

An invalid interface error appears for the following scenarios:

- Interface Mode 'routed' is invalid. Allowed mode is trunk & access.
- Access port which is already allocated to other network.
- Interface which is not available in the switch.

You can use the following functions:

- Create, deploy, view, edit and delete a port channel, vPC, Straight-through FEX, Active-Active FEX, loopback, and subinterface.



- The following features are unsupported for the brownfield migration of switches with the Cisco NX-OS Release 7.0(3)I4(8b) and 7.0(4)I4(x) images:

- FEX on switches other than Cisco N9300 Series switches and Cisco N9500 Series switches with X9500 line cards
- AA-FEX

For information about the platform support for FEX, refer to your platform and NX-OS documentation to check the feature compatibility.

- To edit interfaces associated with fabric links such as intra-fabric links and inter-fabric links, see [Edit interfaces associated with links](#).
- The **flowcontrol** or **priority-flow-control** config is not supported for HIF ports or PO with HIF ports as members.
- When using the REST API for configurations, make sure to set consistent values for the primary fields and NV pairs fields. For example, a REST API post for a single port channel that has different values in certain fields, such as:

- **ifName:** Port-testing123
- **PO_ID:** Port-channel1000

results in two interfaces being created rather than the intended single interface.

- Create tunnel interfaces for Cisco Cloud Services Router 1000v Series (Cisco CSR 1000v Series).
- Create breakout and unbreakout ports.
- Shut down and bring up interfaces.
- Rediscover ports and view interface configuration history.

- Apply host policies on interfaces and vPCs. For example, int_trunk_host, int_access_host, and so on.
- View interface information such as its admin status, operation status, reason, policy, speed, MTU, mode, VLANs, IP/Prefix, VRF, port channel, and the neighbor of the interface.



The **Neighbor** column provides details of connected switches that are discovered, intent links, and Virtual Machine Manager (VMM) connectivity.

The **Status** column displays the following status of an interface:

- Blue: Pending
- Green: In Sync/Success
- Red: Out-of-Sync/Failed
- Yellow: In Progress
- Grey: Unknown/NA
- If an interface is created out-of-band, you need to perform fabric resync or wait for Config Compliance polling before this interface can be deleted. Otherwise, Config Compliance does not generate the correct diff.

However, you cannot add or edit interfaces for ASR 9000 Series Routers and Arista switches.

You can filter and view information for any of the given fields (such as Device Name).



- Ensure that appropriate configurations are deployed on the Fabric before deploying from the Interfaces option, including proper vPC pair configurations. If you add or edit an interface before configurations are deployed on the Fabric, the configuration may fail on the device.
- Deploy any underlays including vPC Pairing in the fabric before deploying any configurations from the interface manager.

The following table describes the action items available in **Actions** menu drop down list when you select an interface.

Field	Description
Actions	
Create interface	Allows you to add a logical interface such as a port channel, vPC, Straight-through FEX, Active-Active FEX, and loopback.

Field	Description
Edit configuration	Allows you to edit and change policies that are associated with an interface.  Access-admin user role cannot edit interfaces associated with link policy such as inter-fabric link or intra-fabric link for easy fabrics. The user role can edit interfaces for LAN classic and IPFM fabrics. If you select the int_monitor interface policy, you are responsible for the interface configuration and management. Nexus Dashboard does not apply any interface configurations nor perform configuration compliance for interfaces that have the int_monitor policy.
Actions > Configuration	
Create subinterface	Allows you to add a logical subinterface.
Multi-attach	Allows you to attach multiple networks when you add an interface from an interface group in the Interfaces .
Multi-detach	Allows you to detach multiple networks when you remove an interface from an interface group in the Interfaces .
Preview	Allows you to preview the interface configuration.
Deploy	Allows you to deploy or redeploy saved interface configurations.
No Shutdown	Allows you to enable an interface (no shutdown or admin up).
Shutdown	Allows you to shut down the interface.
Breakout	Allows you to breakout an interface.
Un-Breakout	Allows you to unbreakout interfaces that are in breakout state.
Actions > Interface group	
Add	Allows you to add or remove an interface to an interface group.
Remove	Allows option to edit an interface group.
Actions > Maintenance	
Rediscover	Allows you to rediscover or recalculate the compliance status on the selected interfaces.
Show commands	Allows you to display the interface show commands. A show command requires show templates in the template library.
History	Allows you to display the interface deployment history details.
Actions > Bulk actions	

Field	Description
Import	Allows you to import the edited interfaces. The following are the limitations during importing the interfaces: - You are not allowed to import interfaces with these policy templates: - All fabric templates with int_fabric or int_ipfm_fabric int_vpc_peer and int_vpc_leaf_tor_assoc int_freeform templates - You must update the mandatory fields fabric name, serial number, interface name, and policy name. - You are not allowed to import the interfaces with the interface name <i>nve</i> and <i>vlan</i> except <i>int_ipfm_vlan</i> policy. You can import the interface with the <i>int_ipfm_vlan</i> policy. - The allowed MTU range for integer values is between 576 and 9216. - The allowed MTU string values is either <i>default</i> or <i>jumbo</i>. - The fabric name, serial number, and interface name must be unique. - You can only import a single policy type for an interface per .csv file. Importing a .csv file with multiple policy types is not allowed.  There is a server property to set the maximum number of rows that can be imported. By default, the property is 200 for import.
Export	Allows you to export the selected interfaces with multiple types of policies to a .csv file. While there is technically no limit on the number of interfaces to export, the number of interfaces included in each exported .csv file is limited to the number of rows that are displayed on the page. For example, if you select all of the interfaces on this page and you have 50 as the entry in the Rows per page field at the bottom of the page, then only the 50 interfaces displayed on this page are exported to the .csv file.
Normalize	Allows you to apply the same interface policies and its parameters on selected interfaces in one shot.
Actions	
Delete	Allows you to delete a logical interface that is created from the Interfaces screen. An interface having a policy that is attached from an overlay and underlay cannot be deleted.

You can disable deployments, or freeze, a fabric in Nexus Dashboard as a network administrator. However, you cannot perform all actions when you freeze the fabric or if the fabric is in monitor mode.

The following table describes the actions you can perform when you freeze a fabric and when you enable the monitor mode for a fabric.

Operations	Nexus Dashboard Mode	
	Freeze Mode	Monitor Mode
Add	Save, Preview	Blocked
Breakout	Blocked	Blocked
Unbreakout	Blocked	Blocked
Edit	Save, Preview	Blocked
Delete	Save, Preview	Blocked
Shutdown	Save, Preview	Blocked
No Shutdown	Save, Preview	Blocked
Show	Supported	Supported
Rediscover	Supported	Supported
Deploy	Blocked	Blocked
Import	Supported	Supported
Export	Supported	Supported

The buttons for the associated operations are grayed out accordingly.

If you perform admin operations (shutdown/no shutdown) on an SVI, which is part of a config profile, successive **Save & Deploy** operations generate the **no interface vlan** command.

For an SVI with no policy, on performing admin operation, that is, shutdown/no shutdown command pushed from **Interface Manager**, the **int_vlan_admin_state** policy is associated with the SVI.

For example, create and deploy the SVI from **switch_freeform**.

```
interface vlan1234
  description test
  no shutdown
  no ip redirects
  no ipv6 redirects
```

If you shutdown the SVI from the interface manager, the **int_vlan_admin_state** policy is associated with the SVI.

Pending diff is shown as:

```
interface Vlan1234
  shutdown
  no ip redirects
  no ipv6 redirects
  description test
  no shutdown
```

Remove the **no shutdown** CLI from the free-form config.

If you have performed an admin operation on the SVI, the device has an interface as a running config. Therefore, post network detach **interface vlan** is still present and the interface is discovered. You need to manually delete the interface from the **Interface Manager**.

The following table describes the fields that appear on the **Manage > Fabrics > Overview > Connectivity > Interfaces** tab.

Field	Description
Interface	Specifies the interface name.
Switch	Specifies the switch name.
Admin status	Specifies the administrative status of the interface. The status can be either Up or Down.
Operational status	Specifies the operational status of the interface. The status can be either Up or Down.
Neighbor	Specifies the interface neighbor.
Reason	Specifies the reason.
Overlay network	Specifies the overlay network.
Sync status	Specifies the sync status. Specifies if the interface status is In-Sync or Out-Of-Sync.
Anomaly level	Specifies the anomaly level.
Interface group	Specifies the interface group to which the interface belongs to.
Port channel ID	Specified the port channel ID.
vPC ID	Specifies the vPC ID.
Speed	Specifies the interface speed.
MTU	Specifies the MTU size.
Discovered configuration mode	Specifies the discovered interface configuration mode of the switch.
VLANs	Specifies the VLANs.
IP/Prefix	Specifies the interface IP/prefix.
VRF	Specifies virtual routing and forwarding instances (VRFs).
Description	Specifies the interface description.  There is a known issue where the description entry is truncated to 64 characters. If the interface description is more than 64 characters, run the snmp ifmib ifalias long command on the switch to increase the description entry length to 256 characters.
Intended configuration mode	Specifies the intended interface configuration mode of the switch.
Policy	Specifies the policy name.

Add interfaces

You can create different types of interfaces, such as Port Channel, Virtual Port Channel (VPC), Loopback, Switch Virtual Interface (SVI), Ethernet, and sub-interfaces—based on your network requirements.

Before you begin

To create a vPC interface, first configure the vPC switches in the fabric topology and deploy the global vPC pairing and peer-link configurations using the Save & Deploy option. The vPC pair then appears in the Select a vPC pair drop-down list when you create the vPC interface.

Follow these steps to add interfaces.

1. [Navigate to the Connectivity page](#).
2. Click the **Interfaces** tab.
3. From the **Actions** drop-down list, choose **Create interface**.

The **Create interface** page appears.

4. From the Type drop-down list, choose the interface type.

Valid values are **Port Channel**, **Virtual Port Channel (VPC)**, **Loopback**, **Switch Virtual Interface (SVI)**, **Ethernet**, and **Tunnel**. The respective interface ID field appears based on the interface type you choose.



When you create a port channel through Nexus Dashboard, add interfaces of the same speed. A port channel that uses interfaces of different speeds does not come up. For example, a port channel with two 10-Gigabit Ethernet ports is valid, but a port channel that combines a 10-Gigabit Ethernet port with a 25-Gigabit Ethernet port is not valid.

5. (If applicable) From the **Mode** drop-down list, choose the mode for the interface.

For example, choose **Dot1q Tunnel** to configure a Dot1q tunnel mode vPC. NOTE: The **Mode** field is not available for **Loopback**, **Switch Virtual Interface (SVI)**, and **Tunnel** interface types.

6. From the **Select a device** drop-down list, choose a device.

For **Virtual Port Channel (VPC)** interface type, use the **Select a vPC pair** drop-down list to choose a vPC switch pair. Devices are listed by fabric, switch, and interface type. To narrow your search, enter the first few characters of the device name in the drop-down list. After you choose the device or vPC pair, Nexus Dashboard automatically populates the interface **ID** and **Policy** fields based on the interface type.



- o To override the auto-populated ID, enter a new value in the interface ID field. Nexus Dashboard applies the new value only if it is available in the Resource Manager pool; otherwise, an error appears.
- o For Cisco Catalyst 9000 Series switches, the loopback ID range for user-defined loopback interfaces is 0-2147483646. For Cisco N9000 Series switches, the loopback ID range for user-defined loopback interfaces is

0-1023.

- o The default policy is **Host**. For **Loopback** and **Switch Virtual Interface (SVI)** interface types, the default policy is **Base**.
- o To apply a custom policy template, choose **User defined** from the **Policy** drop-down list. Then, in the **Template** field, click **Select template** and choose the user-defined policy template you want to apply.

7. Enter values in the required fields under **Policy Options**.

The fields vary according to the interface type you chose.



- o To mirror the configuration of Peer-1 on Peer-2 while creating a vPC, check the Enable Config Mirroring check box. Nexus Dashboard disables the Peer-2 fields and copies the values you enter in the Peer-1 fields to the Peer-2 fields.
- o A trunk port carries untagged packets along with tagged packets. When you assign a default port VLAN ID to a trunk port, all untagged traffic travels on that VLAN, which is known as the native VLAN ID. The native VLAN ID is the VLAN that carries untagged traffic on trunk ports.

8. Click **Save** to save the configurations.

9. (Optional) Click **Preview** to review the configuration before you deploy it.

10. Click **Deploy** to deploy the logical interface.

The newly added interface appears in the **Interfaces** table.

11. To view configurations for the new interface, click the corresponding policy of the new interface from the **Policy** column in the **Interfaces** table. The **Expected config** dialog box opens and displays the CLI configuration that Nexus Dashboard deploys to the switch.

Expected config



MS1 > MS1-LEAF2

interface vPC3

```
1 interface port-channel3
2 switchport
3 switchport mode trunk
4 switchport trunk allowed vlan none
5 mtu 9216
6 vpc 3
7 spanning-tree bpduguard enable
8 spanning-tree port type edge trunk
9 description e1/1-3
10 no shutdown
11
```

MS1 > MS1-LEAF1

interface vPC3

```
1 interface port-channel3
2 switchport
3 switchport mode trunk
4 switchport trunk allowed vlan none
5 mtu 9216
6 vpc 3
7 spanning-tree bpduguard enable
8 spanning-tree port type edge trunk
9 description e1/1-5
10 no shutdown
11
```

Breakout

To breakout an interface, from Nexus Dashboard:

1. [Navigate to the Connectivity page.](#)
2. Click the **Interfaces** tab.
3. Choose the appropriate interface from the list.
4. On the **Interfaces** page, click **Actions > Configuration > Breakout**.

The **Breakout Interfaces** page appears.

5. Choose the required option on the page and click **Breakout**.

The available options are 10g-4x, 25g-4x, 50g-2x, 50g-4x, 100g-2x, 100g-4x, 200g-2x, and Unbreakout.

UnBreakout

You can unbreakout interfaces that are in a breakout state.

1. [Navigate to the Connectivity page.](#)
2. Click the **Interfaces** tab.
3. On the **Interfaces** page, click **Actions > Configuration > UnBreakout**.



The unbreakout option is grayed out for interfaces that are not in breakout state.

Edit interfaces

The **Edit interface(s)** page allows you to configure interfaces, modify applied policies, and manage interface assignments for port channels or virtual port channels (vPCs).

Follow these steps to edit interfaces.

1. [Navigate to the Connectivity page.](#)
2. Click the **Interfaces** tab.
3. Choose the interface that you want to edit.

You can choose multiple interfaces at the same time to edit them sequentially.



You cannot edit multiple port channels and vPCs. You cannot edit interfaces of different types at the same time.

4. From the **Actions** drop-down list, choose **Edit configuration**.

The **Edit interface(s)** page displays parameters based on the template and the policy that you chose.



Nexus Dashboard does not support interface description configuration for FEX

fabric ports with fabric resource templates. While the FEX fabric ports and FEX are provisioned correctly, the description configured through Nexus Dashboard is not displayed on the leaf switch interfaces, the fabric port channel, or the APIC interface configuration wizard.

- a. Optional: Click the policy name, choose a different policy, and then click **Select**. The choice of policies depends on the interface type. Ensure that you choose an appropriate policy for the interface.
- b. If you are editing a trunk host, you can edit these fields under the **General Parameters** tab.
 - **Enable BPDU Guard**—Enables or disables the spanning tree BPDU guard. If enabled, it prevents BPDU traffic.
 - **Configure BPDU Filter**—Enables or disables the spanning tree BPDU filter. If enabled, the interface cannot send nor receive BPDUs.
 - **Spanning-tree Link-type**—Specifies the link type for the spanning tree protocol (STP) to use.
 - **Enable Port Type Fast**—Enables or disables the spanning tree edge port behavior.
 - **MTU**—Maximum transmission unit (MTU) size of the interface.
 - **SPEED**—Interface speed.
 - **Trunk Allowed Vlans**—Specifies the VLANs allowed on the trunk.
 - **Native Vlan**—Specifies the native VLAN.
 - **Interface Description**—Add a description for the interface.
 - **Enable Auto-Negotiation**—Enables or disables auto-negotiation. Auto-negotiation is an optional function of the IEEE 802.3u Fast Ethernet standard that enables devices to automatically exchange information over a link about speed and duplex abilities.
 - **Enable CDP**—Enables or disables Cisco Discovery Protocol (CDP) on the interface.
 - **Enable vPC Orphan Port**—Enables or disables vPC orphan port functionality. When enabled, you can configure the interface as a vPC orphan port, and the secondary peer suspends it during vPC failures.
 - **Port Duplex Mode**—Configures the communication mode of the port to operate in either full-duplex or half-duplex mode.
 - **Bandwidth in kilobits**—Specifies the allowed bandwidth in kilobits of the interface.
 - **Inherit Bandwidth in kilobits**—Specifies the allowed bandwidth in kilobits of all sub-interfaces of this interface.
 - **Debounce Timer**—Specifies the link debounce timer in milliseconds, which is the wait time for the interface when there is a flap in the link. After the specified time passes, Nexus Dashboard checks the link status again to re-confirm the event. If the link is okay at this time, then the interface remains up.
 - **Debounce Link-up Timer**—Specifies the link debounce link-up timer, which is the wait time before bringing up the interface after it was taken down due to a flap.

- **Enable Error Detection**—Enables or disables error detection for access control list. When enabled, if the software on the switch detects an error situation on the interface, the software shuts down that interface and no traffic is sent nor received on that interface.
- **Forwarding Error Correction**—Specifies the Forwarding Error Correction (FEC) mode. Forwarding Error Correction is a method of obtaining error control in data transmission over an unreliable or noisy channel in which the source (transmitter) encodes the data in a redundant way using an Error Correcting Code, and the destination (receiver) recognizes it and corrects the errors without requiring a retransmission.
- **Freeform Config** – Allows custom configurations.
- **Enable Interface** – Enables or disables the interface.
- **Enable Netflow** – Enables or disables Netflow monitoring. Netflow monitoring is supported only if it is enabled on fabric.
- **Netflow Monitor** – Configures Netflow monitoring settings.
- **Netflow Sampler** – Configures Netflow sampling settings.
- **Flow Monitor Configuration**—Configures one or more flow monitor entries on supported IOS XE interfaces. For each entry, choose the IP version, enter the flow monitor name, and choose the traffic direction.
- **Device Tracking Policy**—Applies a device tracking policy to the supported IOS XE interface.



The **Flow Monitor Configuration** and **Device Tracking Policy** fields are available for IOS XE trunk, access, port channel trunk, and port channel access interfaces on Cisco Catalyst switches in Campus VXLAN EVPN fabrics and IOS XE switches in External fabrics. These fields allow the interface intent in Nexus Dashboard to match telemetry-related configuration pushed by Catalyst Center. The referenced flow monitor and device tracking policy must already exist on the switch. If the referenced flow monitor or device tracking policy is not defined, deployment can fail.

- **Enable priority flow control** – Enables or disables priority flow control.
- **Enable QoS Configuration** – Enables or disables quality of service (QoS) settings. When enabled, you can configure a QoS policy specific to the interface. If AI queuing is active on the fabric, Nexus Dashboard applies the **QOS_CLASSIFICATION** policy by default. You must provide a custom policy in the **Custom QoS Policy** field to override this default behavior.
- **Custom QoS Policy** – Applies a custom QoS policy.
- **Custom Queuing Policy** – Applies a custom queuing policy.
- **Switchport monitor** – Enables or disables switchport monitoring.
- **Enable VLAN Mapping** – Enables or disables VLAN mapping.

The feature is applicable to these trunk interface types.

- Ethernet
- Port channel

- Virtual port channel (vPC)

For more information, see [Enable VLAN mapping](#).

c. If you are editing a trunk host, you can edit these fields under the **Storm Control** tab.

- **Configure Traffic Storm Control**—Enables or disables storm control on the interface. Storm control prevents excessive broadcast, multicast, or unicast traffic from disrupting the network by monitoring traffic levels and applying thresholds.
- **Storm Control Action**—Specifies the action to take when a traffic storm is detected. You can see these options.
 - **Shutdown**—Disables the port when a storm is detected.
 - **Trap**—Sends an SNMP trap notification when a storm is detected.
 - **No**—Returns to default settings.
- **Broadcast Storm Control Level in Percentage**—Specifies the broadcast traffic utilization threshold at which point storm control begins to manage broadcast traffic.
- **Multicast Storm Control Level in Percentage**—Specifies the multicast traffic utilization threshold at which point storm control begins to manage broadcast traffic.
- **Unicast Storm Control Level in Percentage**—Specifies the unicast traffic utilization threshold at which point storm control begins to manage broadcast traffic.
- **Broadcast Storm Control Level in PPS**—Specifies the threshold level for broadcast traffic in packets per second (pps). The interface blocks all traffic when traffic utilization exceeds this level.
- **Multicast Storm Control Level in PPS**—Specifies the threshold level for multicast traffic in packets per second. The interface blocks all traffic when traffic utilization exceeds this level.
- **Unicast Storm Control Level in PPS**—Specifies the threshold level for unicast traffic in packets per second. The interface blocks all traffic when traffic utilization exceeds this level.

d. Edit the other policy options fields as necessary.

e. Click **Save**, and then click **Deploy**.

Additional information

This information applies for editing interfaces.

- In a vPC setup, the two switches are in the order the switch names are displayed on the **Edit** page. For example, if **Switch Name** is displayed as *LEAF1:LEAF2*, then Leaf1 is peer switch one and Leaf2 is peer switch two.
- During overlay network deployment on switches, the network can be associated with trunk interfaces. The **Interface** tab displays the trunk interface to network association. You can update such interfaces.

Enable VLAN mapping

You can configure VLAN mapping on trunk interfaces to enhance network management. To enable VLAN mapping, configure the switch to map customer VLANs to provider VLANs. This configuration segregates traffic and ensures proper handling of VLAN IDs across networks.

Follow these steps to configure VLAN mapping on a trunk interface.

1. [Navigate to the Connectivity page.](#)
2. Click the **Interfaces** tab.
3. Choose the interface that you want to edit.

You can choose multiple interfaces at the same time to edit them sequentially.

4. In the **Edit interface(s)** page, check the **Enable VLAN Mapping** check box.

The **Vlan Mapping Entries** table is enabled.

Edit interface(s)

☒ Enable VLAN Mapping

Vlan Mapping Entries*

Filter by attributes

Actions

☐	Customer Vlan	Selective dot1q-tunnel	Customer Inner Vlan	Provider Vlan

No rows found

The **Vlan Mapping Entries** table displays these information.

- o **Customer Vlan** – Represents the Vlan ID used within the customer’s network.
- o **Selective dot1q-tunnel** – Allows selective tunneling of specific Vlans by mapping a range of customer Vlans to provider Vlan.



Selective dot1q-tunnel and port Vlan mapping cannot be present on the same interface.

- o **Customer Inner Vlan** – Refers to the VLAN tag located inside the header of an incoming packet from the customer.
- o **Provider Vlan** – Represents the Vlan ID assigned by the service provider to encapsulate customer traffic for transport across the provider's network.

5. Add a new mapping entry.

Follow these steps to add a new mapping entry.

- a. From the **Actions** drop-down list, click **Add**.

The **Add Item** dialog box opens.

- b. In the **Customer Vlan** field, provide the customer Vlan ID or a range of Vlan IDs.
- c. Check the **Selective dot1q-tunnel** check box, if needed.



Ensure that you configure the **Selective dot1q-tunnel** option to all the customer Vlans. You cannot configure both selective dot1q-tunnel and port Vlan mapping on the same interface.

- d. Optional: In the **Customer Inner Vlan** field, provide the inner Vlan ID from the customer network.



You cannot provide **Customer Inner Vlan** if you choose the **Selective dot1q-tunnel** check box. When you choose **Selective dot1q-tunnel**, Nexus Dashboard uses provider Vlan for tunneling and does not require customer inner Vlan.

- e. In the **Provider Vlan** field, enter the provider Vlan ID to map the customer Vlan.
- f. Click **Save** to save the configuration.

Limitations

- For PVLAN interfaces, you can associate interfaces with the networks only for access and trunk port types.
- For interface policies that are not created from the **Manage > Inventory > Interfaces > Interfaces** page, you can edit some configurations but not change the policy itself. The policy and fields that cannot be edited are grayed out.
- These are some examples of policies that you cannot edit.
 - Loopback interface policies. You use the `int_fabric_loopback` policy to create a loopback interface. You can edit the loopback IP address and description, but not the `int_fabric_loopback` policy instance.

You cannot edit the loopback IP addresses for loopback interfaces that are created automatically while creating and attaching the VRF instances.

- Fabric underlay network interface policies, such as `int_fabric_num`, and fabric overlay network interface (NVE) policies.
- Policies associated with port channels and member ports of port channels, including the port channels and member ports associated with a vPC.
- SVIs created during network and VRF instance creation. The associated VLANs appear in the interfaces list. `!include-catalyst-center-assurance-fields:`

Edit interfaces associated with links

There are two types of links, namely intra-fabric links and inter-fabric links. As the name implies, intra-fabric links are set up between devices within the same VXLAN fabric and are typically used for spine-leaf connectivity. Inter-fabric links are set up between a VXLAN fabric, and typically other External or VXLAN fabrics. They are used for external WAN and/or DCI connectivity. A policy is

associated with each link that effectively states the configuration that is applied to both ends of the link. In other words, the link policy becomes the parent of the individual child interface policies that are associated with the two interfaces that form the link. In this scenario, you must edit the link policy to edit the interface policy fields such as description, IP address, and any per interface freeform configuration.

Follow these steps to edit the interfaces associated with links:

1. [Navigate to the Connectivity page](#).
2. Click the **Interfaces** tab.
3. Choose a link and click **Actions > Maintenance > Rediscover**.

Delete interfaces

To delete the interfaces from Nexus Dashboard, perform the following steps:



- This option allows you to delete only logical ports, port channels, and vPCs. You can delete the interface if it does not have overlay or underlay policy attached.
- When a port channel or vPC is removed, the corresponding member ports get the default policy associated. The default policy can be configured in `server.properties` file.

1. [Navigate to the Connectivity page](#).
2. Click the **Interfaces** tab.
3. Choose the interfaces and click **Actions > Delete**.

You cannot delete logical interfaces created in the fabric underlay.

4. Click **Save**.
5. Click **Deploy** to delete the interface.

Shut down and bring up interfaces

To shut down and bring up the interfaces from Nexus Dashboard:

1. [Navigate to the Connectivity page](#).
2. Click the **Interfaces** tab.
3. Choose the interfaces that you want to shut down or bring up.

Note: For all interfaces except VLANs, you cannot perform a **Shutdown** or **No Shutdown** on interfaces that do not have a policy attached.

4. Click **Actions > Configuration > Shutdown** to disable the selected interfaces. For example, you may want to isolate a host from the network or a host that is not active in the network.

A confirmation page appears where you can save, preview, and deploy the changes.

5. Click **Save** to preview or deploy the changes.

6. Click **Actions > Configuration > No Shutdown** to bring up the selected interfaces.

A confirmation page appears where you can save, preview, and deploy the changes.

7. Click **Save** to preview or deploy the changes.

View interface configuration

To view the interface configuration commands and execute them from Nexus Dashboard, perform the following steps:

1. [Navigate to the Connectivity page](#).
2. Click the **Interfaces** tab.
3. Choose the interface with configurations that you want to view and choose **Actions > Maintenance > Show Commands**.
4. On the **Interface show commands** page, select the action from the **Commands** drop-down list and click **Execute**.

The interface configurations are displayed on the right of the screen.

For show commands, you must have corresponding **show** templates for interface or interface sub types like port channel or vPC, defined in the **Templates**.

Rediscover interfaces

To rediscover the interfaces from Nexus Dashboard:

1. [Navigate to the Connectivity page](#).
2. Click the **Interfaces** tab.
3. Choose the interfaces that you want to rediscover and choose **Actions > Maintenance > Rediscover** to rediscover the selected interfaces.

For example, after you edit or enable an interface, you can rediscover the interface.

View interface history

To view the interface history from Nexus Dashboard, perform the following steps:

1. [Navigate to the Connectivity page](#).
2. Click the **Interfaces** tab.
3. Select the interface and choose **Actions > Maintenance > History** to view the configuration history on the interface.
4. Click **Status** to view each command that is configured for that configuration instance.

Deploy interface configurations

To deploy the interface configuration from Nexus Dashboard, perform the following steps:

1. [Navigate to the Connectivity page](#).
2. Click the **Interfaces** tab.
3. Choose an interface that you want to deploy and choose **Actions > Configuration > Deploy** to deploy or redeploy configurations that are saved for the interface.



You can select multiple interfaces and deploy pending configurations.

After you deploy the interface configuration, the interface status information is updated. However, the overall switch-level state may be in the pending state, which is in blue. The overall switch-level state goes to the pending state whenever there is a change in intent from any module, such as interface, link, policy template update, top-down, or so on. In the pending state, a switch may have pending configurations or switch-level recomputation. The switch-level recomputation occurs when:

- You deploy for the switch
- During a deploy
- During hourly sync

Create external fabric interfaces

You can add and edit port channel, vPC, subinterface, and loopback interfaces for external fabric devices. You cannot add Straight-through FEX and Active-Active FEX functions.

The Breakout port function is only supported for the Cisco N9000, Cisco Nexus 3000, and Cisco Nexus 7000 Series Switches in the external fabric.

When you add an interface to an external fabric device, the Resource Manager is not in sync with the device. So, ensure that the value populated in the ID field (Port-channel ID, vPC ID, Loopback ID, etc) is not previously configured on the switch.

If you want to configure a portchannel in the external fabric, you should add and deploy the **feature_lacp** policy on the switches where the portchannel will be configured.

When an external fabric is set to **Fabric Monitor Mode Only**, you cannot deploy configurations on its switches. If you click **Save & Deploy** in the fabric topology screen, it displays an error message. However, the following settings (available when you right-click the switch icon) are allowed:

- vPC pairing—You can designate a vPC switch pair, but it is only for reference.
- View/edit policy—You can add a policy but you cannot deploy it on the switch.
- Manage interfaces—You can only create intent for adding interfaces. If you try to deploy, edit, or delete interfaces, it results in an error message.

Sync up out-of-band switch interface configurations

Any interface level configuration made outside of Nexus Dashboard (via CLI) can be synced to Nexus Dashboard and then managed from Nexus Dashboard. Also, the vPC pair configurations are automatically detected and paired. This applies to the External and Classic LAN fabrics only. The vPC pairing is performed with the **vpc_pair** policy.



When Nexus Dashboard is managing switches, ensure that all configuration changes are initiated from Nexus Dashboard and avoid making changes directly on the switch.

When the interface config is synced up to the Nexus Dashboard intent, the switch configs are considered as the reference, that is, at the end of the sync up, the Nexus Dashboard intent reflects what is present on the switch. If there were any undeployed intent on Nexus Dashboard for those interfaces before the resync operation, they will be lost.

Guidelines

- Supported in fabrics using the following fabric templates: Data Center VXLAN EVPN, External, and Classic LAN.
- Supported for Cisco Nexus switches only.
- Supported for interfaces that do not have any fabric underlay related policy associated with them prior to the resync. For example, IFC interfaces and intra fabric links are not subjected to resync.
- The time taken by host port resync depends on the number of switches/interfaces to be synchronized.
- Supported for interfaces that do not have any custom policy (policy template that is not shipped with Cisco Nexus Dashboard) associated with them prior to resync.
- Supported for interfaces where the intent is not exclusively owned by a Cisco Nexus Dashboard feature and/or application prior to resync.
- Supported on switches that do not have Interface Groups associated with them.
- Interface mode (switchport to routed, trunk to access, and so on) changes are not supported with overlays attached to that interface.

The sync up functionality is supported for the following interface modes and policies:

Interface Mode	Policies
trunk (standalone, po, and vPC PO)	• int_trunk_host • int_port_channel_trunk_host • int_vpc_trunk_host
access (standalone, po, and vPC PO)	• int_access_host • int_port_channel_access_host • int_vpc_access_host
dot1q-tunnel	• int_dot1q_tunnel_host • int_port_channel_dot1q_tunnel_host • int_vpc_dot1q_tunnel_host
routed	int_routed_host
loopback	int_freeform
sub-interface	int_subif

FEX (ST, AA)	▪ int_port_channel_fex ▪ int_port_channel_aa_fex
breakout	interface_breakout
nve	int_freeform (only in External_Fabric/Classic LAN)
SVI	int_freeform (only in External_Fabric/Classic LAN)
mgmt0	int_mgmt

The interface resync automatically updates the network overlay attachments based on the access VLAN or allowed VLANs on the interface.

After Nexus Dashboard completes the resync operation, the switch interface intent can be managed using normal Nexus Dashboard procedures.

Sync up switch interface configurations

We recommend that you deploy all switch configurations from Nexus Dashboard. In some scenarios, it may be necessary to make changes to the switch interface configuration out-of-band. This will cause configuration drift causing switches to be reported Out-of-Sync.

Nexus Dashboard supports syncing up the out-of-band interface configuration changes back into its intent.

Guidelines and limitations

The following limitations are applicable after syncing up the switch interface configurations to Nexus Dashboard:

- This feature is not supported on ToR/Access switches, or on leaf switches with ToR pairing present.
- The port channel membership changes (once the policy exists) are not supported.
- Changing the interface mode (trunk to access and so on) that have overlays attached is not supported.
- Resync for interfaces that belong to **Interface Groups** are not supported.
- The vPC pairing in **External_Fabric** and **Classic LAN** templates must be updated with the **vpc_pair** policy.
- The resync can be performed for a set of switches and repeated as desired.
- The time taken by host port resync depends on the number of switches/interfaces to be synchronized.
- In **Data Center VXLAN EVPN** fabrics, VXLAN overlay interface attachments are performed automatically based on the allowed VLANs.

Before you begin

- We recommend taking a fabric backup before attempting the interface resync.
- In **External Fabric** and **Classic LAN** fabrics, for the vPC pairing to work correctly, both the

switches must be in the fabric and must be functional.

- Ensure that the switches are **In-Sync** and switch mode must not be **Migration** or **Maintenance**.
- From the **Actions** drop-list, choose **Discovery > Rediscover** to ensure that Nexus Dashboard is aware of any new interfaces and other changes.

Follow these steps to sync up switch interface configurations.

1. Choose **Manage > Fabrics** and single-click on a fabric.

The **Fabric Overview** page appears.

2. Click **Inventory > Switches** and ensure that switches are present in the fabric.
3. Click **Inventory > vPC Pairs** and ensure that vPC pairings are completed.
4. Click **Configuration Policies > Policies** tab and choose one or more switches where the interface intent resync is needed.



- If a pair of switches is already paired with either **no_policy** or **vpc_pair**, select only one switch of the pair.
- If a pair of switches is not paired, then select both the switches.

5. From the **Actions** drop-down list, choose **Add policy**.

The **Create Policy** page appears.

6. On the **Create Policy** page, choose **host_port_resync** from the **Policy** drop-down list.
7. Click **Save**.
8. Click **Inventory > Switches**, then check the **Mode** column for the switches to ensure that they report **Migration**. For a vPC pair, both switches are in the **Migration-mode**.
 - After this step, the switches in the **Topology view** are in **Migration-mode**.
 - Both the switches in a vPC pair are in the migration mode even if one of the switches is placed into this mode.
 - If switches are unintentionally put into the resync mode, they can be moved back to the normal mode by identifying the **host_port_resync** policy instance and deleting it from the **Policies** tab.
9. After the configuration changes are ready to sync up to Nexus Dashboard, navigate to the **Switches** tab and select the required switches.
10. Click **Actions > Recalculate and deploy** to start the resync process.



This process might take some time to complete based on the size of the switch configuration and the number of switches involved.

11. The **Deploy Configuration** page is displayed if no errors are detected during the resync operation. The interface intent is updated in Nexus Dashboard.



If the External_Fabric or Classic LAN fabric is in **Monitored Mode**, an error message indicating that the fabric is in the read-only mode is displayed. This error message can be ignored and doesn't mean that the resync process has

failed.

Close the **Deploy Configuration** page, and you can see that the switches are automatically moved out of the **Migration-mode**. Switches in a vPC pair that were not paired or paired with **no_policy** show up as paired and associated with the **vpc_pair** policy.



The **host_port_resync** policy that was created for the switch is automatically deleted after the resync process is completed successfully.

Interface groups

An interface group consists of multiple interfaces with same attributes. You can create an interface group that allows grouping of host-facing interfaces at a fabric level. Specifically, you can create an interface group for physical Ethernet interfaces, Layer 2 port-channels, and vPCs. You can attach or detach multiple overlay networks to the interfaces in an interface group.

Shared Policy

You can create and add a shared policy for an interface group. This policy ensures update of appropriate configurations for all the interfaces in the Interface group. In the shared policy, all the interfaces will have the same underlay and overlay attributes. When you change the configuration in the shared policy, then that configuration is applied to all the interfaces.

You can see the details of shared policy under the **Policies** column.

Custom policy can also be created by selecting the policy from the template list and **Duplicate Template** to add the additional information. The shared policy must contain the tags **interface_edit_policy**, **interface_edit_shared_policy**, and **int_trunk**.

Guidelines

- Interface groups are only supported for the following fabric template types:
 - **Campus VXLAN EVPN**
 - **Data Center VXLAN EVPN**
 - **Enhanced Classic LAN**
- An interface group is specific to a fabric. For example, consider two fabrics: Fab1 and Fab 2. The interface group IG1 in Fab1 isn't applicable to Fab 2.
- An interface group can only have interfaces of a certain type. For example, you need three separate interface groups if you want to group three types of interfaces such as IG1 for physical Ethernet trunk interfaces, IG2 for Layer 2 trunk port-channels, and IG3 for vPC host trunk ports.
- An interface group can also be created using preprovisioned interfaces.
- Interface groups are supported only to switches with leaf or border roles. For Border Gateway roles, Interface Groups are supported only on vPC BGWs but not on Anycast BGW, BGW Spine, or BGW SuperSpine.
- You can include Layer 2 ToR interfaces in the interface groups.
- Interfaces added to an interface group with a shared policy replaces the individual policy and get the shared policy.
- You can change description and status of each interface in interface group.
- Interface removed from an interface group with a shared policy will set to a default policy.
- VMs should have the same configuration for all the interfaces under the shared policy.
- Shared policy is supported only for **Ethernet** interface type in interface group.
- Ethernet interface groups now have a common policy.
- **Port-Channel** and **vPC** interface types are not supported for adding shared policy in interface group.

- When the MTU value in the shared policy has to be changed, make sure to update the fabric settings with the same value across all switches of that fabric.
- When the **ptp**, **ttag**, and **ttag-strip** option from the shared policy has to be used, make sure to enable PTP globally in fabric settings.
- When the netflow option has to be used, make sure all the interfaces of interface groups are capable of netflow configuration and that it is enabled globally in fabric settings.
- For Layer 2 port-channels and vPCs that are part of an interface group, they can't be deleted until they are de-associated from the interface group even if there are no networks associated with the interface group. Similarly, a trunk port that has no overlay networks but is part of an IG can't be converted to an access port. In other words, you can't change policies for interfaces that are part of an interface group. However, you can edit certain fields for policies.
- For L4-L7 services configuration on leaf switches, trunk ports that are used for services attachment can't be part of interface groups.
- When you perform a per fabric backup of an easy fabric, if there are interface groups created in that fabric, all the associated interface group state is backed up.
- If an easy fabric contains an interface group, then this fabric can't be imported into the MSO. Similarly, if an easy fabric has been added to the MSO, you can't create interface groups for interfaces that belong to switches in the easy fabric.
- The **Add to interface group** and **Remove from interface group** button is enabled only for Admin and Stager users. For all other users, this button is disabled.
- The **Interface Groups** button is disabled in the following circumstances:
 - Select any other interface apart from vPC, Port-channel, and Ethernet.
 - If the interface has a policy attached from another source, for example:
 - If the interface is member of a port-channel or vPC.
 - If the port-channel is member of vPC.
 - If the interface has a policy from underlay or links.



If you select different types of interfaces, the **Interface Group** button is enabled. However, when you try to create or save different types of interfaces to an interface group, an error is displayed.

Create an interface group

To create an interface group from Nexus Dashboard:

1. [Navigate to the Connectivity page](#).
2. Click the **Interface Groups** tab.
3. Choose **Actions > Create interface group**.
4. In the **Create interface group** window, provide an interface group name in the **Interface Group Name** field, choose an Interface Type, and click **Save**.

An interface group name can have a maximum length of 64 characters.



An interface can be added to single interface group only.

5. Click on the **Policy** field.

The **Select Policy** window appears.

6. Choose **int_shared_trunk_host** policy and then click **Select**.

You can add a shared policy to the interface group which can be shared by the interfaces existing in that group. Shared Policy is optional, for upgrades, all the existing interface group will not have a policy.



The policy field supports the **Ethernet** interface type only.

7. Enter the mandatory parameters in the text field and click **Save**.
8. In the **Interfaces** tab, select the interfaces to group and choose **Actions > Add to interface Group**.

The **Add new interface Group** window appears.

9. To create a custom interface group, enter an interface group name in the **Select Interface Group** field and click **Create custom**.

If you have already created an interface group, select it from the **Select Interface Group** drop-down list. Also, if an interface is already part of an interface group, you can move it to a different interface group by selecting the new group from the **Select Interface Group** drop-down list.

You can create interface groups from either the **Interfaces Groups** window or the **Interfaces** window under Fabric Overview.

10. Click **Save**.

In the **Interfaces** window, you can see the interface group name under the **Interface Group** column.

11. To edit an interface group, choose **Actions > Edit Interface Group**. You can update the policy options after you assign the shared policy.



You cannot edit or delete the shared policy template.

Remove interfaces from an interface group

To remove interfaces from an interface group from Nexus Dashboard:

1. [Navigate to the Connectivity page](#).
2. Click the **Interface Groups** tab.
3. Select the interfaces to disassociate from an interface group and choose **Actions > Remove from interface Group**.

A dialog box appears asking whether you want to clear all the associated interfaces.

4. Click **Yes** to proceed.

Note that if there are any networks attached to these interfaces, they are detached as well when you click **Clear**.

Attach networks to an interface group

To attach networks to an interface group from the Nexus Dashboard Web UI:

1. Click on the fabric to launch **Fabric Overview**.
2. Click **Segmentation and Security > Networks**.
3. On the **Networks** tab, select the networks that you need to attach to an interface group and click **Interface Group**.



An overlay network can belong to multiple interface groups. You can select only the networks with a VLAN ID. Otherwise, an error message is displayed.]

4. In the **Interface Groups** window, you can perform the following:
 - o Select an existing interface group from the **Select Interface Group** drop-down list and click **Save**.

For example, you select three networks and the interface group **test**, and click the **Save** button, the following operations are performed in the background:

- a. Nexus Dashboard retrieves interfaces that are part of the interface group **test**.
- b. Nexus Dashboard determines that three networks are added to the interface group **test**. Therefore, it autoattaches these networks to all the interfaces that are part of the interface group **test**.
- c. For each interface, Nexus Dashboard pushes the "switchport trunk allowed vlan add xxxx" command three times for each selected network.



Nexus Dashboard ensures that there's no duplicate configuration intent.

If you click **Clear**, Nexus Dashboard pushes "switchport trunk allowed vlan remove xxx" config intent.

- o Create a custom interface group by entering an interface group name in the **Select Interface Group** field and click **Create new interface group**. Click **Save**.

If you choose this option, make sure to add interfaces to this Interface Group in the **Interfaces** window. As a result, Nexus Dashboard performs the following operations:

- Removes all existing overlay networks that don't belong to the interface group from these interfaces.
- Adds new overlay networks to these interfaces that are part of the interface group but not yet attached to these interfaces.

For more information about associating interfaces to interface groups, see [Create an interface group](#).

5. Click **Actions > Recalculate and deploy** to deploy the selected networks on the switches.

Detach a network from an interface group

This procedure shows how to detach a network from an interface group in the Networks window. Also, you can detach networks when you remove an interface from an interface group in the **Interfaces** window. For more information, see *Removing Interfaces from an Interface Group*.

1. Click on the fabric to launch **Fabric Overview**.
2. Click **Segmentation and Security > Networks**.
3. On the **Networks** tab, select the networks that you need to detach to an interface group and click **Add to interface group**.
4. In the **Add to interface groups** window, select the interface group from the **Select Interface Group** drop-down list and click **Clear** to detach a network.
5. (Optional) Navigate to **Manage > Connectivity > Interfaces**.

Under the **Overlay Network** column, you can see the detached network in the red color for the corresponding interface. Click the network to view the expected config that is struck through.

6. Navigate to **Segmentation and Security > Networks** and choose **Actions > Recalculate and deploy**.

Delete an interface group

An interface group is automatically deleted when it's not in use. You can perform an explicit delete by clicking on **Interface Groups > Actions > Delete interface group**. This check is performed whenever you click the **Clear** button in the **Edit Interface Group** window. There may be exception scenarios where you need to clean up the interface groups explicitly.

For example, assume that you create an interface group called **storageIG** and you add an interface to it. Later, you want to change the interface mapping to another group (**diskIG**). In this case, you would go through these steps:

1. [Navigate to the Connectivity page](#).
2. Click the **Interfaces** tab.
3. Choose the interfaces that you want to remove from the **storageIG** interface group.
4. Click **Actions > Remove from interface group**.
5. With the same interfaces chosen, click **Actions > Add to interface group** and add the interfaces to the **diskIG** interface group.
6. Click the **Interface Groups** tab.
7. Assuming that the **storageIG** interface group no longer has any interfaces because they were moved to the **diskIG** interface group, choose the **storageIG** interface group and click **Actions > Delete interface group**.

Links

You can add links between border switches of different fabrics (inter-fabric links) or between switches in the same fabric (intra-fabric links). You can only create an inter-fabric connection (IFC) for a switch that is managed by Nexus Dashboard.

There are scenarios where you might want to define links between switches before connecting them physically. The links could be inter-fabric or intra-fabric links. Doing so, you can express and represent your intent to add links. The links with intent are displayed in a different color till they are actually converted to functional links. Once you physically connect the links, they are displayed as connected.

Management links might show up in the fabric topology as red colored links. To remove such links, right-click the link and click **Delete Link**.

The Border Spine and Border Gateway Spine roles are added to switch roles for border switches.

You can create links between existing and pre-provisioned devices as well by selecting the pre-provisioned device as the destination device.

These are the subtabs available under **Links**:

- [Links](#)
- [Protocol view](#)

Links

The following table describes the fields that appear on the **Links** tab.

Field	Description
Fabric name	Specifies the name of the Fabric.
Name	Specifies the name of the link. The list of previously created links is displayed. The list contains intra-fabric links, which are between switches within a fabric, and inter-fabric links, which are between border switches in this fabric and switches in other fabrics.
Policy	Specifies the link policy.
Link state	Provides more link state, such as Neighbor Present , Link Present , and so on.
Admin status	Displays the administrative status of the link.
Operational status	Displays the operational status of the link.

The following table describes the action items, in the **Actions** menu drop-down list, that appear on **Fabric Overview > Links > Links**.

Action Item	Description
-------------	-------------

Create link	Allows you to create the following links: • Create inter-fabric links • Create intra-fabric links
Edit	Allows you to edit the selected fabric.
Delete	Allows you to delete the selected fabric.
Import	You can import a CSV file containing details of links to add new links to the fabric. The CSV file should have the following details of links: link template, source fabric, destination fabric, source device, destination device, source switch name, destination switch name, source interface, destination interface, and nvPairs.  You cannot update existing links. The Import Links icon is disabled for an external fabric.
Export	Select Actions > Export to export the links in a CSV file. The following details of links are exported: link template, source fabric, destination fabric, source device, destination device, source switch name, destination switch name, source interface, destination interface, and nvPairs. The nvPairs field consists JSON object.

Protocol view

This tab displays the protocols for the links in the selected fabric.

The following table describes the fields that appear on **Protocol View** tab.

Field	Description
Fabric Name	Specifies the name of the fabric.
Name	Specifies the name of the link.
Is Present	Specifies if the link is present.
Link Type	Specifies the type of link.
Link State	Specifies the state of link.
UpTime	Specifies the time duration from when the link was up.

To delete a protocol for links in a selected fabric, choose the protocol and click **Actions > Delete**.

Create intra-fabric links

To create intra-fabric links:

1. [Navigate to the Connectivity page](#).
2. Click the **Links** tab.
3. Click **Actions > Create**.

The **Link Management - Create Link** page appears.

4. From the **Link Type** drop-down list, choose **Intra-Fabric** since you are creating an IFC. The screen changes correspondingly.

The fields are:

- **Link type**—Choose **Intra-Fabric** to create a link between two switches in a fabric.
- **Link subtype**—This field populates the fabric indicating that this is a link within the fabric.
- **Policy**—You can choose any of the following link policies.
 - **Numbered**-- If the link is between two ethernet interfaces assigned with IP addresses, choose **Numbered**.
 - **Unnumbered**-- If the link is between two IP unnumbered interfaces, choose **Unnumbered**.
 - **vPC peer-keepalive**-- If the link is a vPC peer keep-alive link, choose **vPC peer-keepalive**.
 - **Pre-provision**-- If the link is between two pre-provisioned devices, choose **Pre-provision**. After you click Save & Deploy, an IP address is picked from the underlay subnet IP pool.

Correspondingly, the **Link profile** section field is updated.

- **Source fabric**—The fabric name populates this field since the source fabric is known.
- **Destination fabric**—Choose the destination fabric. For an intra-fabric link, source and destination fabrics are the same.
- **Source device** and **Source interface**—Choose the source device and interface.
- **Destination device** and **Destination interface**—Choose the destination device and interface.



Select the pre-provisioned device as the destination device if you are creating a link between an existing device and a pre-provisioned device.

In the **General** tab in the **Link Profile** section:

- **Interface VRF**—Name of a non-default VRF for this interface.
- **Source IP** and **Destination IP**—Specify the source and destination IP addresses of the source and destination interfaces, respectively.



The **Source IP** and **Destination IP** fields do not appear if you choose the `int_pre_provision_intra_fabric_link` template.

- **Enable interface**—Check or uncheck the check box to enable or disable the **admin** state of the interface.



The **Interface Admin State** field has been renamed to **Enable interface**.

- **MTU**—Specify the maximum transmission unit (MTU) through the two interfaces.
- **Source Interface Description and Destination Interface Description**—Describe the links for later use. For example, if the link is between a leaf switch and a route reflector device, you can enter the information in these fields (link from leaf switch to RR 1 and link from RR 1 to leaf switch). This description gets converted into a configuration, but will not be pushed into the switch. After **Save & Deploy**, it is reflected in the running configuration.
- **Disable BFD Echo on Source Interface** and **Disable BFD Echo on Destination Interface**—Check the check box to disable BFD echo packets on the source and the destination interface.

Note that the BFD echo fields are applicable only when you have enabled BFD in the fabric settings.

- **Source Interface Freeform CLIs and Destination Interface Freeform CLIs**—Enter the freeform configurations specific to the source and destination interfaces. You should add the configurations as displayed in the running configuration of the switch, without indentation. For more information, see [Enabling Freeform Configurations on Fabric Switches](#).

5. Click **Save** at the bottom right part of the page.

You can see that the IFC is created and displayed in the list of links.

6. On the **Fabric Overview Actions** drop-down list, choose **Recalculate Config**.

The **Deploy Configuration** page appears.

It displays the configuration status on the switches. You can also view the pending configurations by clicking the respective link in the **Pending Config** column. It lists the pending configurations on the switch. The **Side-by-Side Comparison** tab displays the running configuration and the expected configuration side-by-side.

Close the **Pending Config** page.

7. From the **Fabric Overview Actions** drop-down list, click **Deploy Config**.

The pending configurations are deployed.

After ensuring that the progress is 100% in all the rows, click **Close** at the bottom part of the page. The **Links** page displays again. In the fabric topology, you can see that the link between the two devices is displayed.

Create inter-fabric links

1. [Navigate to the Connectivity page](#).
2. Click the **Links** tab.



In external fabrics, inter-fabric links support BGW, Border Leaf/Spine, and edge router switches. To create inter-fabric links, perform the following steps:

3. Click **Actions > Create**.

The **Link Management - Create Link** page appears.

- From the **Link type** drop-down box, choose **Inter-fabric** since you are creating an IFC. The page changes correspondingly.

The fields for inter-fabric link creation are as follows:

Field	Description
Link type	Choose Inter-fabric to create an inter-fabric connection between two fabrics, over their border switches.
Link subtype	This field populates the IFC type. Choose the link subtype from the drop-down list: ▪ VRF_LITE ▪ MULTISITE_UNDERLAY ▪ MULTISITE_OVERLAY ▪ L2_DCI For information about VXLAN MPLS interconnection, see Editing Data Center VXLAN EVPN Fabric Settings. For information about routed fabric interconnection, see the section "Creating Inter-Fabric Links Between a Routed Fabric and an External Fabric" in Managing BGP-Based Routed Fabrics.
Policy	The link policy is populated. The policies are autopopulated with corresponding prepackaged default templates that are based on your selection.  You can add, edit, or delete user-defined templates. See Managing Your Template Library for more details.
Source fabric	This field is prepopulated with the source fabric name.
Destination fabric	Choose the destination fabric from this drop-down box.
Source device and Source interface	Choose the source device and Ethernet interface that connects to the destination device.
Destination device and Destination interface	Choose the destination device and Ethernet interface that connects to the source device. Based on the selection of the source device and source interface, the destination information is autopopulated based on Cisco Discovery Protocol information, if available. There is an extra validation that is performed to ensure that the destination external device is indeed part of the destination fabric.

- Complete the necessary configurations in the **General Parameters** tab and any other tabs.

The configuration tabs vary depending on the choices that you make earlier in these procedures.

6. Click **Save** at the bottom-right part of the screen.

You can see that the IFC is created and displayed in the list of links.

7. On the **Fabric Overview > Actions** drop-down list, select **Recalculate Config**.

The **Deploy Configuration** page displays.

It displays the configuration status on the switches. You can also view the pending configurations by clicking the respective link in the **Pending Config** column. It lists the pending configurations on the switch. The Side-by-side Comparison tab displays the running configuration and expected configuration side by side.

Close the **Pending Config** screen.

8. From the **Fabric Overview > Actions** drop-down list, select **Deploy Config**.

The pending configurations are deployed.

After ensuring that the progress is 100% in all the rows, click **Close** at the bottom part of the page. The **Links** page comes up again. In the fabric topology, you can see that the link between the two devices is displayed.

If the two fabrics are member fabrics of an MSD, then you can see the link in the MSD topology too.

What's next:

When you enable the VRF-Lite function using the ToExternalOnly method or using a VXLAN fabric group for inter-fabric connectivity, IFCs are automatically created between the (VXLAN fabric) border/BGW device and connected (external fabric) edge router or core device. When you remove the ER/core/border/BGW device, the corresponding IFCs (link PTIs) to/from that switch are deleted on Nexus Dashboard. Next, Nexus Dashboard removes the corresponding IFC configurations, if any, from the remaining devices on the next **Save & Deploy** operation. Also, if you want to remove a device that has an IFC and overlay extensions over those IFCs, you should undeploy all the overlay extensions corresponding to those IFCs for switch delete to be possible.

1. To undeploy VRF extensions, select the VXLAN fabric and the extended VRFs, and undeploy the VRFs on the VRF deployment page.
2. To delete the IFCs, delete the IFCs from the **Links** tab.
3. Ensure that the fabric switch names are unique. If you deploy VRF extensions on switches with the same name, it leads to an erroneous configuration.

The new fabric is created, the fabric switches are discovered in Nexus Dashboard, the underlay networks that are provisioned on those switches, and the configurations between Nexus Dashboard and the switches are synced.

The remaining tasks are:

- o Provision interface configurations such as vPCs, loopback interface, and subinterface configurations. See [Interfaces](#) for more information.
- o Create overlay networks and VRFs and deploy them on the switches. Refer to [Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric](#).

Establishing inter-fabric connectivity using VRF Lite

External connectivity from data centers is a prime requirement where workloads that are part of a data center fabric can communicate with an outside fabric over WAN or backbone services. To enable Layer 3 for north-south traffic flow, use virtual routing and forwarding instances (VRF)-Lite peering between data center border devices and the external fabric edge routers.

A VXLAN (Virtual Extensible Local Area Network) EVPN (Ethernet Virtual Private Network) based data center fabrics provide connectivity by distributing IP-MAC reachability information among various devices within the fabric. The VRF Lite feature is used for connecting the fabric to an external Layer 3 domain. This can be a border router or a Border Gateway router.

You can enable VRF Lite on the following devices:

- Border
- Border Spine
- Border Gateway
- Border Gateway Spine
- Border Super Spine
- Border Gateway Super Spine

Prerequisites and guidelines

- VRF Lite requires Cisco N9000 Series Cisco Nexus Operating System (NX-OS) Release 7.0(3)I6(2) or later.
- Familiarity with VXLAN BGP EVPN data center fabric architecture and VXLAN overlay provisioning through Nexus Dashboard.
- Fully configured VXLAN BGP EVPN fabrics including underlay and overlay configurations for the various leafs and spine devices, external fabric configuration through Nexus Dashboard, and relevant external fabric device configuration (edge routers, for example).
 - You can configure a VXLAN BGP EVPN fabric (and connectivity to an external Layer 3 domain for north-south traffic flow) manually or using Nexus Dashboard.

This document explains the process to connect the fabric to an edge router (outside the fabric, toward the external fabric) through Nexus Dashboard. You must know how to configure and deploy VXLAN BGP EVPN and external fabrics through Nexus Dashboard.

- You can enable VRF Lite on a physical ethernet interface or on a Layer 3 port channel. A subinterface over a physical interface or Layer 3 port-channel interface is created in Nexus Dashboard at the VRF extension moment for each VRF-Lite link for which the VRF is extended over.
- When you create a VXLAN VRF, ensure that you check the following three fields:

Field	Description
Advertise Host Routes	By default, over the VRF-Lite peering session, only nonhost (/32 or /128) prefixes are advertised. If host routes (/32 or /128) must be enabled and advertised from the border device to the edge/WAN router, check the Advertise Host Routes check box. Route-map does outbound filtering. By default, this check box is disabled.
Advertise Default Route	This field controls whether a network statement 0/0 is enabled under the VRF. This in turn advertises 0/0 route in BGP. By default, this field is enabled. When you choose this check box, this ensures that a 0/0 route is advertised inside the fabric over EVPN route type 5 to the leafs, thereby providing a default route out of the leafs toward the border devices.
Config Static 0/0 Route	By default, this check box is checked. This field controls whether a static 0/0 route to the edge/WAN router must be configured under the VRF on the border device. By default, this field is enabled. If WAN/edge routers are advertising a default route over the VRF-Lite peering to the border device in the fabric, then this field must be disabled. In addition, the Advertise Default Route field must be disabled. The 0/0 route that is advertised over the External Border Gateway Protocol sends over EVPN to the leafs without requiring more configuration. The clean IBGP EVPN separation inside the fabric with eBGP for external out-of-fabric peering provides for this desired behavior.

- To delete a VRF-Lite inter-fabric connections (IFC), remove all the VRF extensions that are enabled on the IFC. Otherwise an error message displays. After you remove the VRF-Lite attachments, recalculate and deploy the fabric to remove any pending Layer 3 extension configurations. Nexus Dashboard removes the per-VRF subinterface and per-VRF External Border Gateway Protocol configuration on the devices.

Sample scenarios

These sections explain different use cases for configuring VRF Lite:

- [Automatic VRF-Lite configuration](#)
- [VRF Lite between a Cisco N9000-based border and a Cisco N9000-based edge device](#)
- [VRF Lite between a Cisco N9000-based border and a non-Cisco device](#)
- [VRF Lite between a Cisco N9000-based border and a non-Nexus device](#)

This is a typical use case of a Cisco ASR 9000-based edge router in managed mode.

Automatic VRF-Lite configuration

Guidelines and limitations for automatic VRF-Lite configuration

- Auto inter-fabric connection (IFC) is supported on Cisco Nexus devices only.
- You can configure Cisco ASR 1000 series routers and Cisco Catalyst 9000 series switches as edge routers. To configure, set up a VRF Lite IFC, and connect it as a border device with easy fabric.

- You can configure Cisco ASR 9000 Series routers as edge routers in managed mode.
- If the device in the External fabric is non-Nexus, you must create IFC manually.
- Ensure that no user policy is enabled on the interface that connects to the edge router. If a policy exists, then the interface will not be configured.



In a brownfield import, configurations that already exist on a switch can be imported into Nexus Dashboard as freeform configurations. If those imported freeform configurations overlap with configurations that Nexus Dashboard expects to generate for an automatic VRF-Lite IFC, the generated configuration preview might show duplicate expected configuration lines for the VRF-Lite link member interface. Review the generated configuration and any imported freeform configurations before deploying.

- Autoconfiguration is supported for the following use cases:
 - **Border** role in the VXLAN fabric and **Edge Router** role in the connected external fabric device
 - **Border Gateway** role in the VXLAN fabric and **Edge Router** role in the connected external fabric device
 - **Border** role to another **Border** role directly



Auto configuration is not provided between two Border Gateways (BGWs).

If VRF Lite is required between other roles, you must deploy it manually on Nexus Dashboard.

- To deploy configurations in the external fabric, you must uncheck the **Fabric Monitor Mode** check box in the external fabric settings. When an external fabric is set to **Fabric Monitor Mode Only**, you cannot deploy configurations on the switches.

VXLAN fabric settings

The following are the two modes in which you can deploy VRF Lite. By default, VRF Lite deployment is set to **Manual**. You can change the settings based on your requirement.

- **Manual**—Use this option to deploy the VRF Lite IFCs manually between the source and the destination devices.
- **Back2Back&ToExternal**—Use this option to automatically configure VRF Lite IFCs between a border switch and the edge or core switches in external fabric or between back-to-back border switches in VXLAN EVPN fabric.



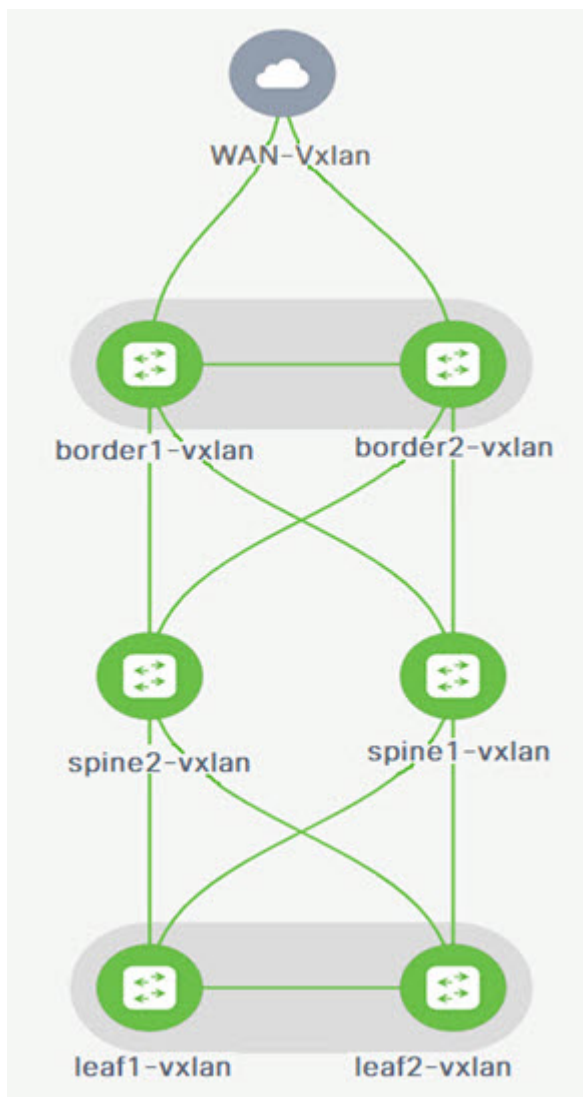
Though VRF Lite mode is set to **Manual** for Nexus Dashboard resource handling, a Data Center Interconnectivity (DCI) subnet is required.

The **Manual** mode is the default mode in VXLAN fabric settings. To change the default mode to another mode, click **Edit fabric settings**. On the **Resource** tab, modify the **VRF Lite Deployment** field to the above mentioned auto configuration mode.

Field	Description
Auto Deploy Peer	This check box is applicable for VRF Lite deployment. When you check this check box, IFCs are automatically created for peer devices. You can check or uncheck this check box when the VRF Lite Deployment field is not set to Manual . The value you choose takes priority. This configuration only affects the new auto-created IFCs and does not affect the existing IFCs.
Auto Deploy Default VRF	When you select this check box, the Auto Generate Configuration on default VRF field is automatically enabled for auto-created VRF Lite IFCs. You can check or uncheck this checkbox when the VRF Lite Deployment field is not set to Manual . The Auto Generate Configuration on default VRF field when set, automatically configures the physical interface for the border device in the default VRF, and establishes an EBGP connection between the border device and the edge device or another border device in a different VXLAN EVPN fabric.
Auto Deploy Default VRF for Peer	When you check this checkbox, the Auto Generate Configuration for NX-OS Peer on default VRF field is automatically enabled for auto-created VRF Lite IFCs. You can check or uncheck this checkbox when the VRF Lite Deployment field is not set to Manual. The Auto Generate Configuration for NX-OS Peer on default VRF field when set, automatically configures the physical interface and the EBGP commands for the peer NX-OS switch.  You can access the Auto Generate Configuration on default VRF and Auto Generate Configuration for NX-OS Peer on default VRF fields for an IFC link by navigating to Actions > Edit > VRF Lite.
Redistribute BGP Route-map Name	Defines the route map for redistributing the BGP routes in the default VRF.
VRF Lite Subnet IP Range	The IP address for VRF Lite IFC deployment is chosen from this range. The default value is 10.33.0.0/16. Ensure that each fabric has its own unique range and is distinct from any underlay range to avoid possible duplication. These addresses are reserved with the Resource Manager.
VRF Lite Subnet Mask	By default, it is set to /30, which is a best practice for point-to-point (P2P) links.

VRF Lite between a Cisco N9000-based border and a Cisco N9000-based edge device

In the following example topology, the DC-VXLAN fabric is on a WAN-VXLAN cloud. The VXLAN fabric has a border leaf role and the WAN-VXLAN cloud has a device with the role edge router. Nexus Dashboard shows the physical and logical representation of the topology with Cisco Discovery Protocol/Link Layer Discovery Protocol (LLDP) link discovery.



In this example, enable a VRF-Lite connection between a DC-VXLAN border leaf and a WAN-VXLAN edge router.

For a VRF-Lite configuration, you must enable External Border Gateway Protocol (EBGP) peering between the fabric's border interfaces and the edge router's interfaces through point-to-point (P2P) connections.

The border physical interfaces are:

- **eth1/1** on **border1-Vxlan**, toward **eth1/1** on **WAN1-Vxlan**.
- **eth1/2** on **border2-Vxlan**, toward **eth1/2** on **WAN1-Vxlan**.

Workflow for configuring an IPv4 or an IPv6 overlay network

Follow these steps to configure an IPv4 or IPv6 overlay network.

1. Create an inter-fabric connection using source and destination IPv4 or IPv6 addresses.

For more information, see [Create inter-fabric links](#).

2. Add a VRF-Lite extension.
3. Attach the VRF and the VRF-Lite extension to the border device.

For more information, see [Attach the VRF and the VRF-Lite extension to the border device](#).

4. Create a network configured with an IPv4 or an IPv6 gateway.

For more information, see the section "Create a network" in [Working with Segmentation and Security for your Nexus Dashboard VXLAN Fabric](#).

5. Add a switch to your network. For more information, see the "Configure switches" section in [Working with Inventory in Your Nexus Dashboard LAN or IPFM Fabrics](#).
6. Recalculate and deploy the configuration.

Create a VRF-Lite inter-fabric link

Follow these steps to create a VRF-Lite inter-fabric link.

1. Navigate to the main **Fabrics** page: **Manage > Fabrics**.
2. In the table showing all of the Nexus Dashboard fabrics that you have already created, locate the row with the Data Center VXLAN EVPN fabric and click on the name of the fabric.

The **Overview** page appears.

3. Navigate to the **Connectivity > Links** tab.
4. On the **Links** tab, click **Actions > Create Link**.

The **Link Management - Create Link** page appears.

5. To configure a VRF-Lite inter-fabric link, enter the required parameters on the **Link Management - Create Link** page.

This table describes the fields on the **Link Management - Create Link** page.

Field	Description
Link Type	Specifies the inter-fabric connection (IFC) link between two different fabrics within Nexus Dashboard. By default, the Inter-Fabric option displays.
Link Sub-Type	Specifies the subtype of link. By default, the VRF_LITE option displays.
Link Template	Specifies the template for the link. The default template for a VRF-Lite IFC is ext_fabric_setup . The template enables the source and destination interfaces as Layer 3 interfaces, configures the no shutdown command, and sets the Maximum Transmission Unit (MTU) value to 9216.
Source Fabric	Nexus Dashboard autodetects and populates the Source Fabric field based on the Cisco Discovery Protocol or LLDP discovery.
Destination Fabric	Nexus Dashboard autodetects and populates the Destination Fabric field based on the Cisco Discovery Protocol or LLDP discovery.
Source Interface	Nexus Dashboard autodetects and populates the Source Interface field based on the Cisco Discovery Protocol or LLDP discovery.

Field	Description
Destination Device	Nexus Dashboard autodetects and populates the Destination Device field based on the Cisco Discovery Protocol or LLDP discovery.
Source Interface	Nexus Dashboard autodetects and populates the Source Interface field based on the Cisco Discovery Protocol or LLDP discovery.
Destination Interface	Nexus Dashboard autodetects and populates the Destination Interface field based on the Cisco Discovery Protocol or LLDP discovery.

6. Enter the necessary field values to configure an inter-fabric link.

The tabs and their fields on the page are explained in the following sections.

- o [General parameters](#)
- o [Advanced](#)
- o [Default VRF](#)
- o [Security](#)

7. When you have completed the necessary configurations, click **Save**.

General parameters

The **General parameters** tab is displayed by default. The fields in this tab are described in the following table.

Field	Description
Source BGP ASN	Specifies the Border Gateway Protocol (BGP) Autonomous System Number (ASN) for the selected source fabric.
Source Address/Mask	IP Specifies the Nexus Dashboard auto allocated IP pool from the resource manager pool of the VRF-Lite subnet pool for the Ethernet1/1 subinterface, the source interface of the IFC. A subinterface is created for each VRF extended over this IFC, and a unique 802.1Q ID is assigned to it. The IP address/mask entered here, along with the BGP neighbor IP field (explained below) is used as the default value for the subinterface that is created as a VRF extension and can be overwritten. For example, an 802.1Q ID of 2 is associated with subinterface Eth 1/1.2 for VRF corporate traffic, and 802.1Q ID of 3 is associated with Eth 1/1.3 and VRF engineering, and so on. The IP prefix is reserved by the Nexus Dashboard resource manager. Ensure that you use a unique IP address prefix for each IFC you create in the topology.

Field	Description
Destination IP Address	Specifies the Nexus Dashboard auto allocated IP pool from the resource manager pool of the VRF-Lite subnet pool. This is a BGP neighbor IP on the device. Inter-fabric traffic from different VRFs for an IFC with the same source IP address (10.33.0.1/30) and destination IP address (10.33.0.2) as an example.
Source Address/Mask IPv6	Specifies the source IPv6 address and mask if you are configuring an IPv6-only inter-fabric link or a dual stack (IPv6 and IPv4) inter-fabric link.
Destination Address IPv6	Specifies the destination IPv6 address if you are configuring an IPv6-only inter-fabric link or a dual stack (IPv6 and IPv4) inter-fabric link.
Destination BGP ASN	Specifies the BGP ASN of the selected destination fabric.
Link MTU	Specifies the default MTU value of 9216.
Auto Generate Configuration for Peer	Specifies to auto generate a VRF-Lite configuration for managed NX-OS neighbor devices. This knob autoconfigures the neighbor VRF on the neighboring managed device. For example, Nexus Dashboard automatically creates a VRF on the edge router inside the WAN-VXLAN External fabric.

What's next: Complete the configurations in another tab if necessary, or click **Save** when you have completed the necessary configurations for this link.

Advanced

The fields in the **Advanced** tab are:

Field	Description
Source Interface Description	Enter a description for the source interface.
Destination Interface Description	Enter a description for the destination interface.
Source Interface Freeform Config	Specifies a freeform configuration for the source interface.
Destination Interface Freeform Config	Specifies a freeform configuration for the destination interface.
Template Configuration for Generation on Peer	Specifies a Python template for VRF-Lite configuration on an NX-OS peer in an External fabric. This is an auto populated field.

What's next: Complete the configurations in another tab if necessary, or click **Save** when you have completed the necessary configurations for this link.

Default VRF

Enter the following details in the **Default VRF** tab.

Field	Description
Auto Configuration default VRF Generate on	Automatically configures the physical interface for the border device for the default VRF. Establishes an External Border Gateway Protocol (EBGP) connection between the border device and its peer in the VRF-LITE IFC (edge router or another border device in a different VXLAN EVPN fabric).
Auto Configuration for NX-OS/IOS XE Peer on default VRF Generate on	Automatically configures the physical interface and the BGP commands for the NX-OS or the IOS XE peer switch corresponding to the default VRF.
Redistribute BGP Route-map Name	Defines the route map used for redistributing the BGP routes in the default VRF.
Default VRF BGP Neighbor Password	Allows you to provide additional security by specifying a password for the BGP neighbor.
Default VRF BGP Password Key Encryption Type	Specifies the encryption type for the BGP password.
VRF Name in NX-OS Peer Switch	Enables you to provide a name for the VRF in the peer switch. By default, the peer device is configured with the default VRF, if left blank.
Enable DCI tracking	Enables Data Center Interconnectivity (DCI) tracking on the physical interface, which allows you to use an existing VRF-Lite link as the VXLAN EVPN multi-site underlay link. Note that you must enable this option only on a border gateway device.
Routing TAG	Specifies the routing tag associated with the interface IP for the default VRF.  All the configuration changes made to the devices as a result of the above configurations are displayed on the Pending Config page.

Security

The fields in the **Security** tab are described in the following table.

For more information on configuring MACsec using QKD, see [Connecting Two Fabrics with MACsec Using QKD](#).

Field	Description
Enable MACsec	Check the check box to enable MACsec on the VRF-Lite inter-fabric connection (IFC) link.  If you enable the Enable MACsec option and you disable the Use Link MACsec Setting option, Nexus Dashboard uses the fabric settings for configuring MACsec on the VRF-Lite inter-fabric link (IFC). When MACsec is configured on the link, Nexus Dashboard generates the following configurations: • Switch-level MACsec configurations if this is the first link that enables MACsec. • MACsec configurations for the link.
Source MACsec Policy/Key-Chain Name Prefix	Specify the prefix for the policy and the key-chain names for the MACsec configuration at the source. The default value is DCI, and you can change the value.
Destination MACsec Policy/Key-Chain Name Prefix	Specify the prefix for the policy and the key-chain names for the MACsec configuration at the destination. The default value is DCI, and you can change the value.
Enable QKD	 If you choose to not enable the Enable QKD option, Nexus Dashboard uses preshared keys provided by the user instead of using the QKD server to generate the keys. If you disable the Enable QKD option, all the fields pertaining to QKD are grayed out.
Use Link MACsec Setting	Check this check box as the override option for using the link settings instead of using the fabric settings.
MACsec Cipher Suite	Choose one of the following MACsec cipher suites for the MACsec policy: • GCM-AES-128 • GCM-AES-256 • GCM-AES-XPB-128 • GCM-AES-XPB-256 The default value is GCM-AES-XPB-256.

Field	Description
MACsec Primary Key String	Specify a Cisco Type 7 encrypted octet string that is used for establishing the primary DCI MACsec session. For AES_256_CMAC, the key string length must be 130 and for AES_128_CMAC, the key string length must be 66. If these values are not specified correctly, an error is displayed when you save the fabric.  The default key lifetime is infinite.
MACsec Primary Cryptographic Algorithm	Choose the cryptographic algorithm used for the primary key string. It can be AES_128_CMAC or AES_256_CMAC. The default value is AES_128_CMAC. You can configure a fallback key on the device to initiate a backup session if the primary session fails.
MACsec Fallback Key String	Specify a Cisco Type 7 encrypted octet string that is used for establishing a fallback MACsec session. For AES_256_CMAC, the key string length must be 130 and for AES_128_CMAC, the key string length must be 66. If these values are not specified correctly, an error is displayed when you save the fabric.  This parameter is mandatory if Enable QKD is not selected.
MACsec Fallback Cryptographic Algorithm	Choose the cryptographic algorithm used for the fallback key string. It can be AES_128_CMAC or AES_256_CMAC. The default value is AES_128_CMAC.
Source QKD Profile Name	Specify the source crypto profile name. The maximum size is 63.
Source KME Server IP	Specify the source IPv4 address for the Key Management Entity (KME) server.
Source KME Server Port Number	Specify the source port number for the KME server.
Source Trustpoint Label	Specify the source authentication trustpoint label. The maximum size is 64.
Destination QKD Profile Name	Specify the destination crypto profile name.
Destination KME Server IP	Specify the destination IPv4 address for the KME server.
Destination KME Server Port Number	Specify the destination port number for the KME server.
Destination Trustpoint Label	Specify the destination authentication type trustpoint label. The maximum size is 64.
Ignore Certificate	Specify if you want to skip verification of incoming certificates.

Click **Save** to save the configuration.

Create a Layer 2 DCI link

Nexus Dashboard added a Layer 2 link template, **ext_l2_dci_link**, for configuring a Layer 2 link for MACsec with QKD. For more information on MACsec with QKD, see [About connecting two NX-OS fabrics with MACsec using QKD](#).

Follow these steps to create a Layer 2 DCI link.

1. On the **Manage > Fabrics** page, click on a Data Center VXLAN EVPN fabric.

The **Overview** page appears.

2. Navigate to the **Connectivity > Links** tab.
3. On the **Links** tab, click **Actions > Create Link**.

The **Link Management - Create Link** page appears.

4. To configure a Layer 2 DCI link, enter the following parameters on the **Link Management - Create Link** page:

Field	Description
Link Type	Choose Inter-Fabric from the drop-down menu.
Link Sub-Type	Choose L2_DCI as the Link Sub-Type .
Link Template	Choose ext_l2_dci_link as the link template for a DCI link.

The link template enables the source and destination interfaces as Layer 2 interfaces.

5. When you have completed the necessary configurations, click **Save**.

Attach the VRF and the VRF-Lite extension to the border device

1. Navigate to the main **Fabrics** page: **Manage > Fabrics**.
2. Click the name of the fabric for which you want to attach the VRF and the VRF-Lite extension.

The **Overview** page displays.

3. Click the **Segmentation and security > VRFs** tab.
4. Choose a **VRF Name** and click **Actions > Edit**.

The **VRF Attachments** page appears.

5. Click the **VRF Attachments** tab.
6. Click **Actions > Edit**.
7. You can edit details in the **Extension** table by clicking the **Edit** link.
8. Toggle the knob to **Attach**.
9. In **Extend**, choose **VRF_LITE** from the drop-down list.
10. In the **Extension** table, choose one switch at a time.

You can edit details in the **Extension** table by clicking the **Edit** link.

The **Edit Extension Details** page displays.

11. In the **Neighbor ASN** field, you can specify an alternate value for the **Neighbor ASN** field. By default, the **Neighbor ASN** field always takes the ASN from the VRF Lite inter-fabric link. Beginning with Nexus Dashboard 4.1.1, you can edit the **Neighbor ASN** field.

If the neighbor is an edge or a core device in an External fabric, Nexus Dashboard uses the **local-as** command in the policy configuration if the Autonomous System Number (ASN) is different from the ASN of the External fabric.

For the pending configurations after editing the **Neighbor ASN** field, see [Pending VRF-Lite extension configurations after modifying a neighbor ASN value](#).

12. Enter the details for **PEER_VRF_NAME**.

This auto deploys the VRF on the neighboring device.

When you extend a VRF-Lite consecutive scenario, the VRF must be in the peer fabric and the VRF name must be the same. If the VRF is not in the peer fabric and if you try to extend VRF Lite, an error message is generated displaying the issue.

When you extend VRF Lite between a VXLAN EVPN fabric and an External fabric, the VRF name can be the same as the name of the source fabric, or the default name, or another VRF name.

13. Enter the required VRF name in the **PEER_VRF_NAME** field.

The child Policy Template Instances (PTIs) for subinterface, VRF creation, and BGP peering on the External fabric have source values that are populated in the PTI. You cannot edit or delete the policies.

14. Follow the above procedure for adding other inter-fabric links.
15. On the **Edit** page, click **Attach-all**, to attach the required VRF extension to the border device.
16. Click **Save**.

Recalculate and deploy the configuration on the VXLAN or the External fabric

1. Navigate to the main **Fabrics** page: **Manage > Fabrics**.
2. Click on the appropriate fabric to navigate to the **Overview** page.
3. Click **Actions > Recalculate and deploy**.
4. Perform the same operation by choosing the required **VRF Name** on the **VRF Attachments** tab and clicking **Actions > Deploy** to initiate the VRF or VRF-Lite configuration on the border device.
5. Alternatively, on the **Overview** page, click **Actions > Recalculate and deploy**.

You can also choose the VRF attachment, edit, and click **Deploy**.

Nexus Dashboard pushes VRF and VRF-Lite configurations to the border devices.

VRF Lite between a Cisco N9000-based border and a non-Cisco device

This section describes the procedure for enabling a VRF-Lite connection between a VXLAN EVPN border leaf device and a non-Cisco device in an External fabric.

We recommend using the meta definition of a device instead of importing devices in an external fabric. This allows VRF-Lite configurations to extend Cisco N9000-managed border devices in VXLAN fabric. Nexus Dashboard does not manage destination non-Cisco devices. You must configure the relevant VRF-Lite configuration on the destination device.

Create new IFC links between a border device and an edge router

Follow the steps documented in [Create inter-fabric links](#) to create new IFC links between a border device and an edge router.

1. On the **Manage > Fabrics** page, click on a Data Center VXLAN EVPN fabric.

The **Overview** page appears.

2. Navigate to the **Connectivity > Links** tab.
3. On the **Links** tab, click **Actions > Create Link**.

The **Link Management - Create Link** page appears.

4. Enter the following required parameters.

Field	Description
Link Type	Specifies the inter-fabric connection (IFC) link between two different fabrics within Nexus Dashboard. By default, the Inter-Fabric option displays.
Link Sub-Type	Specifies the subtype of the link. By default, the VRF_LITE option displays.
Link Template	Specifies the template for the link. The default template for a VRF-Lite IFC is ext_fabric_setup . The template enables the source and destination interfaces as Layer 3 interfaces, configures the no shutdown command, and sets the Maximum Transmission Unit (MTU) value to 9216.
Source Fabric	Select the Source Fabric . This is the easy fabric where a Cisco N9000-based border device resides.
Destination Fabric	Select any External or Classic LAN fabric. The fabric can be in monitor mode as well.
Source Device	Select the Source Device . This is the Cisco N9000-based border device.
Destination Device	You can create a "meta device definition". Type any name and click Save . For example, non-cisco.
Source Interface	Select the interface on the border device where the non-cisco device is connected.

Field	Description
Destination Interface	You can create a "meta device interface". Type any interface name and click Save . For example, gig1, tengig1/10, eth1/1 are the valid interface names.

The tabs and their fields on the page are explained in the following sections.

- [General Parameters](#)
- [Advanced](#)
- [Default VRF](#)

5. When you have completed the necessary configurations, click **Save**.

General Parameters

The **General Parameters** tab is displayed by default. The fields in this tab are described in the following table.

Field	Description
Source BGP ASN	Specifies the Border Gateway Protocol (BGP) Autonomous System Number (ASN) for the selected source fabric.
Source Address/Mask	IP Specifies the IP address and mask for the Ethernet1/5 subinterfaces, the source interface of the IFC. The subinterface is created for each VRF extended over this IFC, and a unique 802.1Q ID is assigned to it. The IP address/mask entered here, along with the BGP Neighbor IP field (explained below) used as the default value for the subinterface that is created as a VRF extension and can be overwritten. For example, an 802.1Q ID of 2 is associated with subinterface Eth 1/5.2 for VRF CORP traffic, and 802.1Q ID of 3 is associated with Eth 1/5.3 and VRF ENG, and so on. The IP prefix is reserved with the Nexus Dashboard resource manager. Ensure that you use a unique IP address prefix for each IFC you create in the topology.
Destination IP Address	Specifies the Nexus Dashboard auto allocated IP pool from the resource manager pool of the VRF-Lite subnet pool. This is a BGP neighbor IP on the device. Interfabric traffic from different VRFs for an IFC with the same source IP address (10.33.0.1/30) and destination IP address (10.33.0.2) as an example.
Source Address/Mask	IPv6 Specifies the source IPv6 address and mask if you are configuring an IPv6-only inter-fabric link or a dual stack (IPv6 and IPv4) inter-fabric link.
Destination Address	IPv6 Specifies the destination IPv6 address if you are configuring an IPv6-only inter-fabric link or a dual stack (IPv6 and IPv4) inter-fabric link.
Destination BGP ASN	Specifies the BGP ASN of the selected destination fabric.
Link MTU	Specifies the default MTU value of 9216.

Field	Description
Auto Deploy Flag	Not applicable as the destination device is a non-Nexus and a non-Cisco device.

Advanced

Enter the following details in the **Advanced** tab. The fields in this tab are described in the following table.

Field	Description
Source Interface Description	Enter a description for the source interface.
Destination Interface Description	Enter a description for the destination interface.
Source Interface Freeform Config	Specifies a freeform configuration for the source interface.
Destination Interface Freeform Config	Specifies a freeform configuration for the destination interface.
Template for Configuration Generation on Peer	Specifies a Python template for VRF-Lite configuration on an NX-OS peer in an external fabric. This is an auto populated field.

What's next: Complete the configurations in another tab if necessary, or click **Save** when you have completed the necessary configurations for this link.

Default VRF

Enter the following details in the **Default VRF** tab. The fields in this tab are described in the following table.

Field	Description
Auto Generate Configuration on default VRF	Automatically configures the physical interface for the border device for the default VRF. Establishes an External Border Gateway Protocol (EBGP) connection between the border device and its peer in the VRF-LITE IFC (edge router or another border device in a different VXLAN EVPN fabric).
Auto Generate Configuration for NX-OS/IOS XE Peer on default VRF	Automatically configures the physical interface and the BGP commands for the NX-OS or the IOS XE peer switch corresponding to the default VRF.
Redistribute BGP Route-map Name	Defines the route map used for redistributing the BGP routes in the default VRF.
Default VRF BGP Neighbor Password	Allows you to provide additional security by specifying a password for the BGP neighbor.
Default VRF BGP Password Encryption Type	Specifies the encryption type for the BGP password.

Field	Description
VRF Name in NX-OS Peer Switch	Enables you to provide a name for the VRF in the peer switch. By default, the peer device is configured with the default VRF, if left blank.
Enable DCI tracking	Enables Data Center Interconnectivity (DCI) tracking on the physical interface, which allows you to use an existing VRF-Lite link as the VXLAN EVPN multi-site underlay link. Note that you must enable this option only on a border gateway device.
Routing TAG	Specifies the routing tag associated with the interface IP for the default VRF.  All the configuration changes made to the devices as a result of the above configurations are displayed on the Pending Config page.

Attach the VRF and the VRF-Lite extension to the border device

Follow these steps to attach the VRF and the VRF-Lite extension to the border device.

1. On the **Manage > Fabrics** page, click on the Data Center VXLAN EVPN fabric.
2. On the **Overview** page, navigate to the **Segmentation and security > VRFs** tab.
3. Click on a **VRF Name**.

The **VRF Attachments** page displays.

4. Click the **VRF Attachments** tab.
5. Choose a **VRF Name** and click **Actions > Edit**.

The **Extension** page displays.

6. Click **Attach-all** to attach the required VRF-Lite extension to the border device and then click **Save**.

Recalculate and deploy the configuration on the VXLAN fabric

1. On the **Manage > Fabrics** page, click on the appropriate fabric to navigate to the **Overview** page.
2. Click **Actions > Recalculate, sync and deploy**.
3. Perform the same action by choosing the required **VRF Name** on the **Segmentation and security > VRF attachments** tab and clicking **Actions > Deploy** to initiate the VRF or VRF-Lite configuration on the border device.

VRF Lite between a Cisco N9000-based border and a non-Nexus device

In this example, you can enable VRF-Lite connections between a DC-VXLAN border leaf and a non-Nexus device in an External fabric.

The following are the supported platforms:

- ASR 9000 and ASR 1000

- Cisco Catalyst 9000
- CSR 1000v
- NCS 5501 and NCS 5001
- Cisco 8000



The ASR 9000 series switch is supported in managed mode with an edge router role.

Configuration compliance is enabled for IOS XR switches in an External fabric, similar to Cisco Nexus switches configured on an External fabric. Nexus Dashboard pushes the configuration at the end of the deployment.



Ensure that the VXLAN BGP EVPN border device is active.

Follow these steps to configure VRF Lite between a Cisco N9000-based border and a non-Nexus device.

1. Navigate to **Manage > Fabrics** to create an External fabric.
2. Choose **Create Fabric**. For information on how to create a fabric, see [Creating Fabrics and Fabric Groups](#).
3. After configuring your External fabric, navigate to the **Manage > Inventory** tab and click **Actions > Add Switches**. For more information, see [Working with Inventory in Your Nexus Dashboard LAN or IPFM Fabrics](#).

You do not need to configure the Simple Network Management Protocol (SNMP) for IOS-XR discovery of switches. Nexus Dashboard uses Secure Shell (SSH) for IOS XR device discovery.

To add non-Nexus devices to External fabrics, see the section "Adding Non-Nexus Devices to External Fabrics" in [External Connectivity Network](#).

4. On the **Add Switches** page, click **Choose Fabric**.
5. Click **IOS-XR** from the **Device Type** drop-down list.
6. After the router is discovered, you can view the switch name in the **Discovery Results** field.
7. Choose the discovered router and add it to the External fabric.

Ensure that the **Discovery Status** displays **OK** in the **Status** column.

An edge router role is supported.

After successful discovery, you can view the links between the devices on the **Links** tab.

8. To create a VRF-Lite Inter-Fabric Connection (IFC) for an External fabric with a Cisco N9000 border leaf, choose the appropriate link and click **Edit** from the **Actions** drop-down list.



You need to perform this step from the VXLAN fabric, not the External fabric.

9. On the **Link Management-Edit Link** page, fill in the required details for creating an inter-fabric connection link. For more information, see [Create inter-fabric links](#).

A few fields are auto-populated.

If the **Auto Generate Configuration for Peer** field is enabled, and the edge device is an IOS XR device, enter **ios_xr_Ext_VRF_Lite_Jython** in the **Template for Configuration Generation on Peer** field. If **Auto Generate Configuration for Peer** is enabled, and the edge device is an IOS XE device, enter **ios_xe_Ext_VRF_Lite_Jython** in the **Template for Configuration Generation on Peer** field.

10. Deploy the configuration on the IOS XR router.

VRF-Lite configurations

Cisco N9000 border device configurations

Border-VXLAN (base border configurations) generated by template **ext_base_border_vrflite_11_1**

```
switch configure terminal
switch(config)#
ip prefix-list default-route seq 5 permit 0.0.0.0/0 le 1
ip prefix-list host-route seq 5 permit 0.0.0.0/0 eq 32
route-map extcon-rmap-filter deny 10
    match ip address prefix-list default-route
route-map extcon-rmap-filter deny 20
    match ip address prefix-list host-route
route-map extcon-rmap-filter permit 1000
route-map extcon-rmap-filter-allow-host deny 10
    match ip address prefix-list default-route
route-map extcon-rmap-filter-allow-host permit 1000
ipv6 prefix-list default-route-v6 seq 5 permit 0::/0
ipv6 prefix-list host-route-v6 seq 5 permit 0::/0 eq 128
route-map extcon-rmap-filter-v6 deny 10
    match ipv6 address prefix-list default-route-v6
route-map extcon-rmap-filter-v6 deny 20
    match ip address prefix-list host-route-v6
route-map extcon-rmap-filter-v6 permit 1000
route-map extcon-rmap-filter-v6-allow-host deny 10
    match ipv6 address prefix-list default-route-v6
route-map extcon-rmap-filter-v6-allow-host permit 1000
```

Border-VXLAN VRF-Lite extension configuration

```
switch configure terminal
vrf context CORP
    ip route 0.0.0.0/0 2.2.2.2
exit
router bgp 100
```

```

vrf CORP
  address-family ipv4 unicast
    network 0.0.0.0/0
  exit
  neighbor 2.2.2.2
    remote-as 200
  address-family ipv4 unicast
    send-community both
    route-map extcon-rmap-filter out
configure terminal
interface ethernet1/1.2
  encapsulation dot1q 2
  mtu 9216
  vrf member CORP
  ip address 2.2.2.22/24
  no shutdown
configure terminal

```

WAN-VXLAN (External fabric edge device) VRF-Lite extension configuration

```

switch configure terminal
vrf context CORP
  address-family ipv4 unicast
  exit
router bgp 200
  vrf CORP
    address-family ipv4 unicast
    neighbor 10.33.0.2
      remote-as 100
    address-family ipv4 unicast
      send-community both
    exit
  exit
  neighbor 10.33.0.6
    remote-as 100
  address-family ipv4 unicast
    send-community both
configure terminal
interface ethernet1/1.2
  mtu 9216
  vrf member CORP
  encapsulation dot1q 2
  ip address 10.33.0.1/30
  no shutdown

```

```
interface ethernet1/2.2
  vrf member CORP
  mtu 9216
  encapsulation dot1q 2
  ip address 10.33.0.5/30
  no shutdown
  configure terminal
```

Pending VRF-Lite extension configurations after modifying a neighbor ASN value

Pending configuration on an edge device

```
vrf context dsvni-1
  address-family ipv4 unicast
exit
router bgp 50000
  vrf dsvni-1
    address-family ipv4 unicast
      neighbor 14.100.0.1
      remote-as 400.4000
      local-as 50012
      address-family ipv4 unicast
        send-community both
  configure terminal
interface ethernet1/8.2
  vrf member dsvni-1
  mtu 9216
  encapsulation dot1q 2
  ip address 14.100.0.2/30
  no shutdown
  configure terminal
```

Pending configuration on a border device

```
router bgp 400.4000
  vrf dsvni-1
    address-family ipv4 unicast
      advertise l2vpn evpn
      redistribute direct route-map fabric-rmap-redist-subnet
      maximum-paths ibgp 2
      network 0.0.0.0/0
    exit
  address-family ipv6 unicast
    advertise l2vpn evpn
```

```
redistribute direct route-map fabric-rmap-redis-subnet
maximum-paths ibgp 2
exit
neighbor 14.100.0.2
remote-as 50012
address-family ipv4 unicast
send-community both
route-map extcon-rmap-filter out
configure terminal
```

About connecting two NX-OS fabrics with MACsec using QKD

Media Access Control Security (MACsec) is the IEEE 802.1AE standard for authenticating and encrypting packets between two MACsec-capable devices.

With this feature, you can connect two fabrics using inter-fabric links with MACsec, either using a quantum key distribution (QKD) server for secure exchange of encryption keys, or by providing a preshared key. If you choose not to enable QKD, Nexus Dashboard configures preshared keys supplied by the user instead of using quantum keys generated by the QKD server. You need to enable MACsec to enable the use of the QKD server for MACsec.



Nexus Dashboard does not manage the QKD server.

Nexus Dashboard provides support for MACsec for inter-fabric links for the following fabric types:

- Data Center VXLAN EVPN
- Enhanced Classic LAN
- External Connectivity Network

Connecting two fabrics with MACsec using QKD

The following sections provide information about connecting two fabrics with MACsec using QKD:

- [Benefits of connecting two fabrics with MACsec using QKD](#)
- [Supported switches and Cisco NX-OS releases for connecting two fabrics with MACsec using QKD](#)
- [Supported configurations for connecting two fabrics with MACsec using QKD](#)
- [Guidelines for connecting two fabrics with MACsec using QKD](#)
- [Limitations for connecting two fabrics with MACsec using QKD](#)
- [Workflow for connecting two fabrics with MACsec using QKD](#)
- [Troubleshooting connecting two fabrics with MACsec using QKD](#)

Benefits of connecting two fabrics with MACsec using QKD

- Supports generation of quantum keys for encryption using a QKD server for protecting the privacy and authenticity of data.

- Instead of using preshared keys between the endpoints, Nexus Dashboard uses the keys from the QKD server that are difficult to break.
- Helps to ensure data confidentiality by providing strong encryption at Layer 2 for VRF-Lite inter-fabric links. For more information, see the section "Create a VRF-Lite Inter-Fabric Link" in [VRF Lite](#).
- Provides a secure connectivity association between fabrics.

Supported switches and Cisco NX-OS releases for connecting two fabrics with MACsec using QKD

For more information about MACsec configuration, which includes supported platforms and releases, see the [Configuring MACsec](#) chapter in the *Cisco N9000 Series NX-OS Security Configuration Guide*.

Supported configurations for connecting two fabrics with MACsec using QKD

You can connect two fabrics with MACsec using QKD for the following use cases:

First Fabric Type	Second Fabric Type
Data Center VXLAN EVPN	Data Center VXLAN EVPN
Data Center VXLAN EVPN	External Connectivity Network
Data Center VXLAN EVPN	Enhanced Classic LAN
Enhanced Classic LAN	Enhanced Classic LAN
Enhanced Classic LAN	External Connectivity Network

Guidelines for connecting two fabrics with MACsec using QKD

- You can change MACsec global parameters in the fabric settings at any time.
- MACsec and CloudSec can coexist on a border gateway (BGW) device.
- MACsec status of a link with MACsec enabled is displayed on the **Links** page.
- Brownfield migration of devices with MACsec configured is supported using switch and interface freeform configurations.

Limitations for connecting two fabrics with MACsec using QKD

This feature does not support the following fabric types:

- BGP
- Campus VXLAN EVPN
- VXLAN EVPN Multi-Site

This feature does not support inter-fabric links between two External Connectivity Network fabrics. You can use a freeform configuration to achieve this use case.

Workflow for connecting two fabrics with MACsec using QKD

1. Before enabling MACsec using QKD on the fabric and the associated switches, you need to do the following:
 - a. Create the RSA keys.

- b. Associate the key pair to the trust point.
- c. Get the signed certificate and upload it to Nexus Dashboard. For more information, see [Upload a certificate for connecting two fabrics with MACsec using QKD](#).
- d. Manually add the signed certificate to the appropriate switches.

For more information, see the section "Configuring CAs and Digital Certificates" in the [Cisco N9000 Series NX-OS Security Configuration Guide](#).

2. Create or edit one of the supported fabric types to enable MACsec and the QKD server.

For more information, see the following sections:

- [Edit fabric settings for a Data Center VXLAN fabric for enabling MACsec using QKD](#)
- [Edit fabric settings for an Enhanced Classic LAN fabric for enabling MACsec using QKD](#)

3. Create a VRF-Lite inter-fabric link or a Layer 2 inter-fabric link.

For more information, see [Working with Connectivity in Your Nexus Dashboard LAN Fabrics](#).

Troubleshooting connecting two fabrics with MACsec using QKD

If you cannot configure MACsec with QKD on the physical interfaces of the link, an error displays when you click **Save**.

The error may occur for the following reasons:

- You are not using the minimum required release version for Cisco NX-OS. For more information, see [Supported switches and Cisco NX-OS releases for connecting two fabrics with MACsec using QKD](#).
- The interface is not MACsec-capable.

Upload a certificate for connecting two fabrics with MACsec using QKD

1. Navigate to **Admin > Certificate Management > CA Certificates**.
2. Click **Add CA Certificate**.

The **Upload Certificate - CA** dialog box displays.

3. Drag and drop your CA certificate file to the dialog box or browse to the location of your certificate file.

The following are the accepted file types:

- .pem
- .cer
- .key
- .crt

4. Click **Save**.

Edit fabric settings for connecting two fabrics with MACsec using QKD

Enable MACsec on the fabric and on each inter-fabric link to configure MACsec using a quantum key distribution server (QKD).

When you enable MACsec on an inter-fabric link, Nexus Dashboard uses the local fabric settings for prepopulating the MACsec parameters. Switches in an External Connectivity Network fabric take MACsec settings from a Data Center VXLAN EVPN fabric or from an Enhanced Classic LAN fabric. An External Connectivity Network fabric, Data Center VXLAN EVPN fabric, or an Enhanced Classic LAN fabric can be either the source or the destination fabric.

Each MACsec-enabled inter-fabric link includes an override option to use the link's local settings instead of from the fabric settings. The override option, **Use Link MACsec Setting**, allows switches in external fabrics to use a different QKD server from the one used in the Data Center VXLAN EVPN fabric or the Enhanced Classic LAN fabric. Nexus Dashboard autopopulates MACsec parameters from the fabric settings if the **Use Link MACsec Setting** option is not enabled.

Edit fabric settings for an Enhanced Classic LAN fabric for enabling MACsec using QKD

1. Choose **Manage > Fabrics**.
2. Create a new fabric or edit an existing one.



If you are editing an Enhanced Classic LAN fabric, you can change MACsec parameters at any time. You need to perform a **Recalculate and deploy** operation to generate new configurations based on the updated fabric settings.

- o If you are creating a new fabric:
 - a. Click **Create Fabric**.
 - b. Click **Create new LAN fabric**, then click **Next**.
 - c. Choose the **Classic LAN** fabric type, then click **Next**.
 - d. Click **Advanced** in the **Configuration mode** field, then configure the necessary settings to create the new fabric.

See [Editing Classic LAN Fabric Settings](#) for more information.

- e. In the **Advanced settings** step in the fabric creation process, click **Security**.

Go to Step 3.

- o If you are editing an existing fabric:
 - a. Choose an existing **Enhanced Classic LAN** fabric.

The **Overview** window for that fabric appears.

- b. Choose **Actions > Edit Fabric Settings**.
- c. Click **Fabric Management**.
- d. Click the **Security** tab.

3. Locate these fields in **Security** and make the necessary configurations:

Field	Description
Enable DCI MACsec	Check the check box to enable MACsec on DCI links.
Enable QKD	Check the check box to enable the QKD server for generating quantum keys for encryption.  If you choose to not enable the Enable QKD option, Nexus Dashboard uses preshared keys provided by the user instead of using the QKD server to generate the keys. If you disable the Enable QKD option, all the fields pertaining to QKD are grayed out.
DCI MACsec Cipher Suite	Choose one of the following DCI MACsec cipher suites for the DCI MACsec policy: • GCM-AES-128 • GCM-AES-256 • GCM-AES-XPN-128 • GCM-AES-XPN-256 The default value is GCM-AES-XPN-256.
DCI MACsec Primary Key String	Specify a Cisco Type 7 encrypted octet string that is used for establishing the primary DCI MACsec session. For AES_256_CMAC, the key string length must be 130 and for AES_128_CMAC, the key string length must be 66. If these values are not specified correctly, an error is displayed when you save the fabric.  The default key lifetime is infinite.
DCI MACsec Primary Cryptographic Algorithm	Choose the cryptographic algorithm used for the primary key string. It can be AES_128_CMAC or AES_256_CMAC. The default value is AES_128_CMAC. You can configure a fallback key on the device to initiate a backup session if the primary session fails.
DCI MACsec Fallback Key String	Specify a Cisco Type 7 encrypted octet string that is used for establishing a fallback MACsec session. For AES_256_CMAC, the key string length must be 130 and for AES_128_CMAC, the key string length must be 66. If these values are not specified correctly, an error is displayed when you save the fabric.  This parameter is mandatory if Enable QKD is not selected.
DCI MACsec Fallback Cryptographic Algorithm	Choose the cryptographic algorithm used for the fallback key string. It can be AES_128_CMAC or AES_256_CMAC. The default value is AES_128_CMAC.

Field	Description
QKD Profile Name	Specify the crypto profile name. The maximum size of the crypto profile is 63.
KME Server IP	Specify the IPv4 address for the Key Management Entity (KME) server.
KME Server Port Number	Specify the port number for the KME server.
Trustpoint Label	Specify the authentication type trustpoint label. The maximum size is 64.
Ignore Certificate	Enable this check box to skip verification of incoming certificates.
MACsec Status Report Timer	Specify the MACsec operational status periodic report timer in minutes.

- Click **Save** to save the fabric settings.
- From the **Actions** drop-down list, choose **Recalculate and deploy** to apply the new MACsec settings to links that have MACsec enabled and where the **Use Link MACsec Setting** option is disabled.

Edit fabric settings for a Data Center VXLAN fabric for enabling MACsec using QKD

- Choose **Manage > Fabrics**.
- Create a new fabric or edit an existing one.



If you are editing an Enhanced Classic LAN fabric, you can change MACsec parameters at any time. You need to perform a **Recalculate and deploy** operation to generate new configurations based on the updated fabric settings.

- o If you are creating a new fabric:
 - Click **Create Fabric**.
 - Click **Create new LAN fabric**, then click **Next**.
 - Choose the **VXLAN** fabric type, then choose the **Data Center VXLAN EVPN** subtype, then click **Next**.
 - Click **Advanced** in the **Configuration mode** field, then configure the necessary settings to create the new fabric.

See [Editing Data Center VXLAN EVPN Fabric Settings](#) for more information.

- e. In the **Advanced settings** step in the fabric creation process, click **Security**.

Go to Step 3.

- o If you are editing an existing fabric:
 - Choose an existing **Enhanced Classic LAN** fabric.

The **Overview** window for that fabric appears.

- b. Choose **Actions > Edit Fabric Settings**.
- c. Click **Fabric Management**.
- d. Click the **Security** tab.

The DCI MACsec parameters are the same as for the **Enhanced Classic LAN** fabric type, as the Data Center VXLAN EVPN fabric type has MACsec parameters for intra-fabric links. For more information, see [Edit fabric settings for an Enhanced Classic LAN fabric for enabling MACsec using QKD](#).

For more information about the remaining fields, see [Editing Data Center VXLAN EVPN Fabric Settings](#).

3. Click **Save** to configure the MACsec with QKD feature on each switch in the fabric.
4. From the **Actions** drop-down list, choose **Recalculate and deploy**.

Edit fabric settings for an External Connectivity Network fabric for enabling MACsec using QKD



With this release, no MACsec parameters are added to the External Connectivity Network fabric.

Switches in an External Connectivity Network fabric take MACsec settings from a Data Center VXLAN EVPN fabric or from an Enhanced Classic LAN fabric.

For more information, see the section "Creating an External Fabric" in [Editing External Fabric Settings](#).

Disable MACsec using QKD

The process described in this section is for disabling MACsec for an inter-fabric link for the supported fabric type. For the supported fabric types, see [About connecting two NX-OS fabrics with MACsec using QKD](#).

The following operations will disable MACsec and QKD from a link:

- Disables MACsec on the inter-fabric link using QKD or a preshared key.
- If this is the last link where MACsec is enabled, Nexus Dashboard deletes the switch-level MACsec configuration from the switch.

1. Navigate to the **Links** window for your fabric.

- a. Navigate to **Manage > Fabrics**.
- b. Click on the fabric where you want to disable MACsec using QKD

The **Overview** window for that fabric appears.

- c. Choose **Connectivity > Links**.
2. Choose the link on which you want to disable MACsec with QKD.
 3. From the **Link Management - Edit Link** page, navigate to the **Security** tab and unselect the following options:
 - **Enable MACsec**

- **Enable QKD**

4. Click **Save**.
5. From the **Fabric Overview > Inventory > Switches** tab, select **Actions > Deploy** to remove the MACsec configuration from the switch.

Routing policies

Introduced in Nexus Dashboard release 4.1.1, routing policies are used to define and control how routes are distributed, filtered, or manipulated between two VXLAN EVPN fabrics or between a VXLAN EVPN fabric and an external fabric in your Nexus Dashboard. These policies give you granular control over routing behavior, enabling the implementation of specific network requirements.

You will use these components to define routing policies in your Nexus Dashboard:

- [L3Outs](#)
- [Route-maps](#)
- [ACLs](#) (access control lists)
- [Prefix-list](#)
- [Community-list](#)
- [Extended community-list](#)

Navigate to the Routing policies page

1. Navigate to the main **Fabrics** page:

Manage > Fabrics

2. In the table showing all of the Nexus Dashboard fabrics that you have already created, locate the VXLAN EVPN or external fabric where you want to configure routing policies.
3. Single-click on that fabric.

The **Overview** page for that fabric appears.

4. Click the **Connectivity** tab.
5. Click the **Routing policies** subtab.

L3Outs

An L3Out (Layer 3 Outside) refers to a logical construct that enables connectivity between two VXLAN EVPN fabrics or between a VXLAN EVPN fabric and an external fabric.

Guidelines and limitations: L3Outs

Supported

L3Outs are supported in these areas:

- On these fabric types:
 - VXLAN EVPN (iBGP)
 - VXLAN EVPN (eBGP)
 - External and Inter-Fabric Connectivity

- o AI VXLAN EVPN

Unsupported

- L3Outs are not supported on these fabric types:
 - o Both fabrics in an L3Out cannot be external fabrics. At least one fabric must be another type.
 - o You cannot create an L3Out from multi-site/fabric group context. You should create the L3Out from one of the member fabrics or from the standalone fabrics.
- L3Outs are not supported with these switch roles:
 - o Leaf
 - o Spine
 - o Super spine
 - o ToR
 - o Access
- When creating an L3Out, in the [2. Reachability configuration](#) step, port channel is not supported with these options:
 - o **Layer-2 trunk**: Not supported on layer 2 port channel
 - o **L3out (Routed)**: Not supported on layer 3 port channelIf you use either of those options in that step, the link that you choose cannot be a port channel.

Working with L3Outs

To view information about L3Outs in your Nexus Dashboard:

1. [Navigate to the Routing policies page](#).
2. Click the **L3Outs** subtab.

Any already-configured L3Outs are displayed with this information:

Field	Description
Name	Shows the name of the L3Out. Click the link in the Name column to display information for this L3Out.
Connected fabric	Shows the destination fabric that is connected in this L3Out. Nexus Dashboard connects two fabrics using an L3Out (for example, you might use an L3Out to connect fabric1 and fabric2). In this example, if you are in fabric1 and you click on the link in the Connected fabric column, the page for fabric2 appears, which is the destination fabric that fabric1 is connected to using this L3Out.
IP version	Shows the IP version used in this L3Out.
Routing interface type	Shows the routing interface type used in this L3Out.
Reachability protocol	Shows the reachability protocol used in this L3Out.
Attach	Shows if this L3Out is attached or not.

Field	Description
Fabric 1 VRF	Shows the VRF associated with fabric 1 in this L3Out.
Fabric 2 VRF	Shows the VRF associated with fabric 2 in this L3Out.
Physical links	Shows the number of physical links used in this L3Out.

3. Configure an L3Out, if necessary.

These options are available to configure L3Outs:

- [Create an L3Out](#)
- [Edit an L3Out](#)
- [Attach an L3Out](#)
- [Detach an L3Out](#)
- [Delete an L3Out](#)

Create an L3Out

To create an L3Out:

1. [Navigate to the Routing policies page.](#)
2. Click the **L3Outs** subtab.
3. Choose **Actions > Create L3Out**.
4. Navigate through the steps in the L3Out wizard:
 - [1. Connect fabrics](#)
 - [2. Reachability configuration](#)
 - [3. Peering and filters](#) (shown if you choose **BGP** in the **Reachability protocol** field in [1. Connect fabrics](#))
 - [3. Fabric 1 routes](#) (shown if you choose **Static Routes** in the **Reachability protocol** field in [1. Connect fabrics](#))
 - [4. Fabric 2 routes](#) (shown if you choose **Static Routes** in the **Reachability protocol** field in [1. Connect fabrics](#))
 - [Summary](#)

1. Connect fabrics

1. Make the necessary configurations in the **Connect fabrics** page:

Field	Description
Name	Enter the name of the L3Out connection.
Fabric 1	Choose the source fabric.
Fabric 2	Choose the destination fabric.

Field	Description
Fabric 1 VRF	Choose the VRF to use for fabric 1 or click + Create VRF to create a new VRF for this fabric. See the "Create a VRF" section in the Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric for more information.
Fabric 2 VRF	Choose the VRF to use for fabric 2 or click + Create VRF to create a new VRF for this fabric. See the "Create a VRF" section in the Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric for more information.
Configure L3Out connections for	Choose whether you want to configure L3Out connections for: ▪ Both fabric 1 and fabric 2, ▪ Fabric 1 only, or ▪ Fabric 2 only. Configuration changes will only be pushed to the switches in the selected fabrics.
IP version	Choose the IP version for this L3Out: ▪ IPv4 ▪ IPv6 ▪ Both
Reachability protocol	Choose the reachability protocol for this L3Out: ▪ BGP ▪ Static routes
Routing interface type	Choose the routing interface type for this L3Out: ▪ Subinterface ▪ Routed ▪ SVI  If you change the interface type, all the previously configured values become invalid as they no longer apply to the new interface type. Therefore, it is essential to review and update the configuration to ensure it accurately reflects the new interface type.
Attach	Toggle the switch to the right to attach the L3Out connection or to the left to unattach the L3Out connection.

2. Click **Next**.

You advance to the [2. Reachability configuration](#) step in the process.

2. Reachability configuration

1. Review the existing reachability configurations displayed in the **Reachability configuration** page, if necessary.

Existing links with the appropriate policy between the two selected fabrics will be displayed here. The link templates used are:

- **VRF lite (Subinterface)** for the subinterface type,
- **Layer-2 trunk** for the SVI type, and
- **L3out (Routed)** for the routed type.

Wherever possible, Nexus Dashboard will attempt to allocate non-conflicting IP addresses, Dot1Q tags or VLAN IDs for these links. If Nexus Dashboard is not able to auto-allocate these configurations, fill in the details manually before proceeding.

Resources are allocated from the following ranges, specified in the fabric settings of VXLAN fabrics.

- **VLAN ID:** VRF VLAN Range
- **Dot1q tag:** Subinterface Dot1q Range
- **IPv4 addresses:** VRF Lite Subnet IP Range / VRF Lite Subnet Mask
- **IPv6 addresses:** VRF Lite IPv6 Subnet Range / VRF Lite IPv6 Subnet Mask



IPv4 and/or IPv6 information will be displayed in the table based on the choice that you made in the **IP version** field.

For each link, the table will display this information:

Field	Description
Fabric 1 switch	Shows the switch in fabric 1 that is part of the link.
Fabric 1 interface	Shows the interface in the fabric 1 switch that is part of the link. If you chose Subinterface in the Routing interface type field, this will specify the name of the subinterface (for example, Ethernet1/1.20).
Fabric 2 switch	Shows the switch in fabric 2 that is part of the link.
Fabric 2 interface	Shows the interface in the fabric 2 switch that is part of the link. If you chose Subinterface in the Routing interface type field, this will specify the name of the subinterface (for example, Ethernet2/1.20).
Fabric 1 IPv4	Specifies the IP address assigned to the inter-fabric connection in fabric 1. This is assigned either to the physical interface, the subinterface, or the SVI based on the choice that you made in the Routing interface type field in 1. Connect fabrics .
Fabric 1 IPv6	The IP subnet should be unique (that is, it should not overlap with existing IP subnets of both the fabrics).

Field	Description
Fabric 2 IPv4	Specifies the IP address assigned to the inter-fabric connection in fabric 2. This is assigned either to the physical interface, the subinterface, or the SVI based on the choice that you made in the Routing interface type field in 1. Connect fabrics .
Fabric 2 IPv6	
IPv4 mask length	Specifies the IPv4 mask or IPv6 prefix length.
IPv6 prefix length	
Dot1q tag	Displayed if you chose Subinterface in the Routing interface type field in 1. Connect fabrics . Specifies the 802.1Q tag assigned to the subinterface for this link (min:2, max:4093).
MTU	Shows the MTU configuration for the link.
PIM	Specifies whether PIM/PIMv6 is enabled on the interface in switch 1. For VXLAN fabrics, ensure that the corresponding fabric settings are enabled when enabling this setting: • Enable IPv4 Tenant Routed Multicast (TRM) • Enable IPv6 Tenant Routed Multicast (TRMv6)
PIMv6	
Fabric 1 interface description	Shows the interface description in the fabric 1 switch that is part of the link.
Fabric 2 interface description	Shows the interface description in the fabric 2 switch that is part of the link.
VLAN ID	Displayed if you chose SVI in the Routing interface type field in 1. Connect fabrics . Specifies the VLAN associated with the SVI in this L3Out.
Switch1 netflow	Specifies whether netflow is enabled on the interface in switch 1. For VXLAN fabrics, ensure that the corresponding fabric settings are enabled when enabling netflow-related options: Enable Netflow
Switch1 netflow monitor	Specifies the netflow monitor for the link in switch 1.
Switch2 netflow	Specifies whether netflow is enabled on the interface in switch 2. For VXLAN fabrics, ensure that the corresponding fabric settings are enabled when enabling netflow-related options: Enable Netflow

Field	Description
Switch2 netflow monitor	Specifies the netflow monitor for the link in switch 2.

2. Determine the action to take for reachability configurations for this L3Out.
 - o If you want to create a new link for this L3Out, go to Step 3.
 - o If you want to edit the configuration of a link in this table, click the box next to that row to choose it and click **Actions > Edit**.
 - o If you want to remove the configuration of a link in this table, click the box next to that row in the table to choose it and click **Actions > Remove link from L3Out**.

3. Determine if you are adding an already-configured link or adding a new physical link.

- o If you are adding an already-configured link:
 - a. Click **Actions > Add link to configure**.
 - b. In the **Add link to configure** page, locate the link that you want to add.

The table in the **Add link to configure** page displays all the discovered links.

- c. Click **Add**.

You are returned to the **Reachability configuration** page with the link displayed in the table. Go to Step 4.

- o If you are creating a new link for this L3Out:
 - a. Click **Actions > Add physical link**.

The **Add physical link** page appears.

- b. In the **Add physical link** page, choose the switches and interfaces that is part of the link:

Field	Description
Fabric 1	
Switch	Choose the switch in fabric 1 that is part of the link.
Interface	Choose the interface in fabric 1 that is part of the link.
Fabric 2	
Switch	Choose the switch in fabric 2 that is part of the link.
Interface	Choose the interface in fabric 2 that is part of the link.
MTU	Enter the MTU configuration for the link.
Speed	Choose the speed for the link.

- c. Click **Save**.

You are returned to the **Reachability configuration** page with the link displayed in the table.



- Any links that you create in the **Add physical link** or **Add link to configure** page will remain as link policies in the system even if you cancel the L3Out creation. You can view, edit, or delete these link policies through **Connectivity > Links**.
- If you want to add a link where one end is a meta switch, then add the link in **Links (Connectivity > Links)**, outside of the L3Out workflow.

4. Complete your configurations in the **Reachability configuration** page.
 - Repeat Step 3 to create additional links for this L3Out, if necessary.
 - When you have finished configuring links for this L3Out, click **Next**.
 - If you chose **BGP** in the **Reachability protocol** field in [1. Connect fabrics](#), you advance to the [3. Peering and filters](#) step in the process.
 - If you chose **Static Routes** in the **Reachability protocol** field in [1. Connect fabrics](#), you advance to the [3. Fabric 1 routes](#) step in the process.

3. Peering and filters

The **Peering and filters** page is displayed only if you chose **BGP** in the **Reachability protocol** field in [1. Connect fabrics](#).



IPv4 and/or IPv6 information will be displayed in the table based on the choice that you made in the **IP version** field.

1. Make the necessary configurations in the **Peering and filters** page:

Field	Description
General	
The following fields provide general configuration options that are applied to both fabrics in this L3Out.	
BGP authentication	Put a check in the box to enable BGP authentication. If you enable this option in the L3Out, you must also enable the Enable BGP Authentication field in the fabric settings in the VXLAN fabrics that are part of this L3Out. See Editing Data Center VXLAN Fabric Settings for more information.
BFD	Put a check in the box to enable the BFD feature. If you enable this option in the L3Out, you must also enable the Enable BFD field in the fabric settings in the VXLAN fabrics that are part of this L3Out. See Editing Data Center VXLAN Fabric Settings for more information.

Field	Description
BGP keep alive timer	Specify the amount of time for the BGP keep alive timer. The value that you enter in the BGP keep alive timer field should be less than the value that you enter in the BGP hold timer field.
BGP hold timer	Specify the amount of time for the BGP hold timer. The value that you enter in the BGP hold timer field should be greater than the value that you enter in the BGP keep alive timer field.
Fabric 1 settings and Fabric 2 settings The following fields are available under the Fabric 1 settings and Fabric 2 settings columns. The descriptions are identical but will apply to either fabric 1 or fabric 2, depending on which column you apply the settings in.	
BGP authentication key type	Choose the BGP authentication key type to use for the fabric: ▪ 3des ▪ Cisco Type6 ▪ Cisco Type7
Encrypted BGP authentication key	Specify the Encrypted BGP authentication key to use for the fabric.
Fabric IPv4/IPv6 route-map in	Choose the IP route-map in to use for the fabric or click + Add route-map to create a new route-map (see Create a route-map for more information).
Fabric IPv4/IPv6 route-map out	Choose the IP route-map out to use for the fabric or click + Add route-map to create a new route-map (see Create a route-map for more information).
Advertise default route (VXLAN fabric only)	Field is enabled by default. Check the box to advertise the default route internally.
Advertise host route (VXLAN fabric only)	Check the box to advertise /32 and /128 host routes to the L3Out BGP peer.
Configure static default route (VXLAN fabric only)	Field is enabled by default. Check the box to enable the flag to automatically generate a static default route configuration.
Soft-reconfiguration	Choose the soft-configuration to use for the fabric: ▪ Inbound always ▪ Inbound ▪ Disabled
default-originate	Check the box to originate a default route towards this peer. This advertises a default route (0.0.0.0/0) to a specific neighbor in the fabric, even if the router itself doesn't have a default route in its routing table.
Local ASN	Choose the local ASN for the fabric if you want the router to appear as a different ASN to the L3Out BGP neighbor.
no-prepend	Check the box to prevent the real ASN from being prepended to the AS-path in the fabric.

Field	Description
replace-as	Available if the no-prepend option is checked above. Check the box to replace the private AS number from outbound updates. This prepends only the local-as number to update to the eBGP neighbor in the fabric.
as-override	Check the box to override the matching AS-number while sending updates. This replaces the neighbor ASN in the AS-path with the local ASN when advertising routes to that neighbor in the fabric.
allows-in	Specify the number of times the local ASN can occur in the AS path. This allows the BGP router to accept routes that contain its own ASN in the AS-path, bypassing the BGP's default loop prevention.
disable-peer-as-check	Check the box to disable the checking of the peer AS-number while advertising. This allows routers in the fabric to advertise routes even when the peer ASN is present in the AS-PATH.
Log neighbor change	Check the box to enable the BGP log neighbor change function.

2. Click **Next**.

Because you chose **BGP** in the **Reachability protocol** field in the [1. Connect fabrics](#) step in the process, you advance to the [Summary](#) step in the process.

3. Fabric 1 routes

The **Fabric 1 routes** page is displayed only if you chose **Static routes** in the **Reachability protocol** field in [1. Connect fabrics](#).

1. Review the existing fabric 1 routes displayed in the **Fabric 1 routes** page, if necessary.

The existing fabric 1 routes are displayed with this information:

Field	Description
Switch	Shows the switch in fabric 1 used in this route.
IP version	Shows the IP version in fabric 1 used in this route.
IP prefix/mask	Shows the IPv4 prefix/mask in fabric 1 used in this route.
Next hop address	Shows the next hop address in fabric 1 used in this route.
Route preference	Shows the route preference in fabric 1 used in this route.
Next hop name	Shows the next hop name in fabric 1 used in this route.
Routing tag	Shows the routing tag in fabric 1 used in this route.
Track object number	Shows the track object number in fabric 1 used in this route. See Configure a track for more information.
Next hop VRF	Shows the next hop VRF in fabric 1 used in this route.

2. Click **Actions > Create** to create a route for fabric 1.

The **Create Fabric 1 routes** page appears.

3. Make the necessary configurations in the **Create Fabric 1 routes** page:

Field	Description
Switch	Choose the switch to use in fabric 1 in this route.
IP version	Choose the IP version to use in fabric 1 in this route. Options are IPv4 , IPv6 , or both .
IPv4 prefix/mask	Field available if you chose IPv4 or both in the IP version field above. Enter the IPv4 prefix/mask to use in fabric 1 in this route.
IPv6 prefix/mask	Field available if you chose IPv6 or both in the IP version field above. Enter the IPv6 prefix/mask to use in fabric 1 in this route.
Next hop address	Enter the next hop address to use in fabric 1 in this route.
Route preference	Enter the route preference to use in fabric 1 in this route.
Next hop name	Enter the next hop name to use in fabric 1 in this route.
Routing tag	Enter the routing tag to use in fabric 1 in this route.
Track object number	Enter the track object number to use in fabric 1 in this route, where the static route is installed in the routing table only if the track is up. See Configure a track for more information before using this option.
Next hop VRF	Enter the next hop VRF to use in fabric 1 in this route.

4. Click **Add**.

You return to the **Fabric 1 routes** page.

- If you want to create another route for fabric 1, click **Actions > Create** and repeat these steps.
- If you are finished creating routes for fabric 1, click **Next**.

You advance to the [4. Fabric 2 routes](#) step in the process.

4. Fabric 2 routes

The **Fabric 2 routes** page is displayed only if you chose **Static routes** in the **Reachability protocol** field in [1. Connect fabrics](#).

1. Review the existing fabric 2 routes displayed in the **Fabric 2 routes** page, if necessary.

The existing fabric 2 routes are displayed with this information:

Field	Description
Switch	Shows the switch in fabric 2 used in this route.
IP version	Shows the IP version in fabric 2 used in this route.
IP prefix/mask	Shows the IPv4 prefix/mask in fabric 2 used in this route.
Next hop address	Shows the next hop address in fabric 2 used in this route.
Route preference	Shows the route preference in fabric 2 used in this route.
Next hop name	Shows the next hop name in fabric 2 used in this route.

Field	Description
Routing tag	Shows the routing tag in fabric 2 used in this route.
Track object number	Shows the track object number in fabric 2 used in this route. See Configure a track for more information.
Next hop VRF	Shows the next hop VRF in fabric 2 used in this route.

- Click **Actions > Create** to create a route for fabric 2.

The **Create Fabric 2 routes** page appears.

- Make the necessary configurations in the **Create Fabric 2 routes** page:

Field	Description
Switch	Choose the switch to use in fabric 2 in this route.
IP version	Choose the IP version to use in fabric 2 in this route. Options are IPv4 , IPv6 , or both .
IPv4 prefix/mask	Field available if you chose IPv4 or both in the IP version field above. Enter the IPv4 prefix/mask to use in fabric 2 in this route.
IPv6 prefix/mask	Field available if you chose IPv6 or both in the IP version field above. Enter the IPv6 prefix/mask to use in fabric 2 in this route.
Next hop address	Enter the next hop address to use in fabric 2 in this route.
Route preference	Enter the route preference to use in fabric 2 in this route.
Next hop name	Enter the next hop name to use in fabric 2 in this route.
Routing tag	Enter the routing tag to use in fabric 2 in this route.
Track object number	Enter the track object number to use in fabric 2 in this route, where the static route is installed in the routing table only if the track is up. See Configure a track for more information before using this option.
Next hop VRF	Enter the next hop VRF to use in fabric 2 in this route.

- Click **Add**.

You return to the **Fabric 2 routes** page.

- If you want to create another route for fabric 2, click **Actions > Create** and repeat these steps.
- If you are finished creating routes for fabric 2, click **Next**.

You advance to the [Summary](#) step in the process.

Summary

- Review the information provided in the **Summary** page and verify that the information is correct for this L3Out.

If any information is incorrect, click **Back** to navigate back to the necessary area to correct that

information.

2. If the information provided in the **Summary** page is correct for this L3Out, click **Save**.

You are returned to the **L3Outs** page, with the newly-configured L3Out listed. Go to [Attach an L3Out](#) to attach the newly-configured L3Out, if necessary.

Edit an L3Out

To edit an L3Out:

1. [Navigate to the Routing policies page](#).
2. Click the **L3Outs** subtab.
3. Choose an L3Out from the configured L3Outs listed in the **L3Outs** page.
4. Choose **Actions > Edit**.
5. Make the necessary changes in the appropriate pages in the configured L3Out, then click **Save**.

Attach an L3Out

You must attach a configured L3Out for the configuration to be generated for the specified fabrics.



For any VRFs that you associated with the L3Out, if you did not attach those VRFs while you configuring the L3Out, attaching the L3Out will automatically attach those VRFs as well.

To attach an L3Out:

1. [Navigate to the Routing policies page](#).
2. Click the **L3Outs** subtab.
3. Choose an L3Out from the configured L3Outs listed in the **L3Outs** page.
4. Choose **Actions > Attach**.

A **Success** *routing-policy-name attached successfully* message appears at the bottom right corner of the page.

Detach an L3Out

You must detach an L3Out before you can delete it.

To detach an L3Out:

1. [Navigate to the Routing policies page](#).
2. Click the **L3Outs** subtab.
3. Choose an L3Out from the configured L3Outs listed in the **L3Outs** page.
4. Choose **Actions > Detach**.

A **Success** *routing-policy-name detached successfully* message appears at the bottom right corner of the page.

Delete an L3Out

You must detach an L3Out before you can delete it. See [Detach an L3Out](#) for more information.

To delete an L3Out:

1. [Navigate to the Routing policies page](#).
2. Click the **L3Outs** subtab.
3. Choose an L3Out from the configured L3Outs listed in the **L3Outs** page.
4. Detach the L3Out using the procedures in [Detach an L3Out](#).
5. Choose **Actions > Delete**.
6. Click **Confirm** in the confirmation message to delete the L3Out.

Configure a track

In order to use the **Track object number** option, you must first configure a track using a freeform configuration. For more information, see "Enable freeform configurations on fabric switches" in [Working with Inventory in Your Nexus Dashboard LAN or IPFM Fabrics](#).

You will see an error message similar to the following if you attempt to push the static route with a track without having configured a track first:

```
C01-BGW-1(config)# ip route 1.1.1.0/24 10.1.1.1 track 1 name test tag 12345
% Track object for object-id 1 associated with the route is not yet created, command
rejected.
```

After configuring a track using a freeform configuration, you should see a message similar to the following, indicating that the operation was successful:

```
DC01-BGW-1(config)# track 1 ip route 10.1.1.1/24 reachability
DC01-BGW-1(config-track)# ip route 1.1.1.0/24 10.1.1.1 track 1 name test tag 12345
DC01-BGW-1(config)#
DC01-BGW-1(config)# sh track 1
Track 1
IP Route 10.1.1.1/24 Reachability
Reachability is DOWN
0 changes, last change never
Tracked by:
IPv4 Static Route 1
```

Route-maps

A route-map is a policy-based tool used to control and manipulate route redistribution, filtering, and policy application within a Nexus Dashboard fabric. Route-maps provide a flexible way to match specific conditions and apply actions to influence the behavior of routing protocols or traffic flows.

To view information about route-maps in your Nexus Dashboard:

1. [Navigate to the Routing policies page.](#)
2. Click the **Route-maps** subtab.

Any already-configured route-maps are displayed with this information:

Field	Description
Name	Shows the name for that route-map.
Type	Shows the type of route-map.
List items	Shows the number of items that this route-map has. Each item under the route-map has an unique sequence number and one or more match/set rules.
Associations	Shows the number of L3Outs that the route-map is associated with.

3. Configure a route-map, if necessary.

These options are available to configure route-maps:

- [Create a route-map](#)
- [Edit a route-map](#)
- [Delete a route-map](#)

Create a route-map

To create a route-map:

1. [Navigate to the Routing policies page.](#)
2. Click the **Route-maps** subtab.
3. Choose **Actions > Create route-map**.
4. In the **Create route-map** page, enter a name for the route-map in the **Name** field.
5. Review the existing route-map items, if necessary.

Any already-configured route-map items are displayed with this information:

Field	Description
Sequence number	Shows the sequence number of the route-map item.
Action	Shows the action assigned to this route-map item.
Route-map entries	Shows the number of route-map entries in this route-map item.

6. In the **Route-map items** area, click **Actions > Add route-map item**.
7. In the **Add route-map item** page, enter the following information:
 - a. In the **Sequence number** field, enter the sequence number.
 - b. In the **Action** field, choose the action:

- **permit:** The route-map permits set operations.
- **deny:** The route-map denies set operations.

c. In the **Route-map entries** area, click **+ Add route-map entry**.

d. In the **Rule type** field, choose the type:

- **Match:** A match rule is a condition-based filter or selector that defines the criteria for identifying traffic or routes in a policy. See [Match rules](#) for the options available under **Match**.
- **Set:** A set rule defines the actions to take when a route or traffic matches specific criteria (defined by a match rule). These actions typically involve modifying attributes such as metrics, next-hop addresses, route tags, or BGP communities. See [Set rules](#) for the options available under **Set**.

Match rules

1. Choose **Match** as the **Rule type**.
2. In the **Rule** area, choose the rule for your match rule.

This table provides the information necessary to configure the different match rules.

Match rule type	Description	Options
match-ipv4-acl	Used to evaluate and match IPv4 routes based on Access Control Lists (ACLs). These ACLs define specific matching criteria for IPv4 addresses or subnets, enabling you to filter or apply actions to routes that meet the specified conditions.	In the Matched IPv4 Access-list Name field, choose an existing ACL or click + Create ACL to create a new one. See Create an ACL for more information.
match-ipv6-acl	Used to evaluate and match IPv6 routes based on Access Control Lists (ACLs). These ACLs define specific matching criteria for IPv6 addresses or subnets, enabling you to filter or apply actions to routes that meet the specified conditions.	In the Matched IPv6 Access-list Name field, choose an existing ACL or click + Create ACL to create a new one. See Create an ACL for more information.
match-ipv4-prefix-list	Used to match IPv4 routes based on a prefix list. A prefix list is a collection of rules that specify which IPv4 address ranges (prefixes) should be matched or filtered. The match-ipv4-prefix-list statement enables you to apply routing policies that act on routes matching these prefixes.	In the Matched IPv4 Prefix-list Names field, choose an existing prefix-list or click + Create prefix-list to create a new one. See Create a prefix-list for more information.
match-ipv6-prefix-list	Used to match IPv6 routes based on a prefix list. A prefix list is a collection of rules that specify which IPv6 address ranges (prefixes) should be matched or filtered. The match-ipv6-prefix-list statement enables you to apply routing policies that act on routes matching these prefixes.	In the Matched IPv6 Prefix-list Names field, choose an existing prefix-list or click + Create prefix-list to create a new one. See Create a prefix-list for more information.

Match rule type	Description	Options
match-community	Used to evaluate and filter routes based on their BGP community attributes. This feature allows you to create routing policies that act on specific routes tagged with one or more BGP communities or patterns of communities.	In the Matched Community-list Name field, choose an existing community-list or click + Create community-list to create a new one. See Create a community-list for more information.
match-extcommunity	Used to evaluate and filter routes based on their BGP extended community attributes. This feature allows you to create routing policies that act on routes tagged with specific BGP extended communities, which are 64-bit values providing additional granularity beyond standard 32-bit BGP communities.	In the Matched Extended community-list Name field, choose an existing extended community-list or click + Create extended community-list to create a new one. See Create a community-list for more information.
match-tag	Used to evaluate and filter routes based on their route tags. A route tag is a numeric identifier applied to routes, often used for marking or categorizing routes for specific policy actions, such as filtering, redistribution, or prioritization.	In the Matched Tag Values field, enter the appropriate match tag values.

3. Click **Save**.

You are returned to the **Add route-map item** page.

4. Make additional configurations in the **Add route-map item** page, if necessary.

- If you want to create another route-map entry, click **+ Add route-map entry**.
- If you want to delete the route-map entry that you've already created, click **Delete** in that route-map entry area.
- If you are finished adding the route-map list item, click **Save**.

You are returned to the **Create route-map** page, with the newly-created route-map list item listed.

5. Make additional configurations in the **Create route-map** page, if necessary.

- If you want to create another route-map item, click **Actions > Add route-map item**.
- If you want to edit an existing route-map item, choose the route-map item and click **Actions > Edit**.
- If you want to delete an existing route-map item, choose the route-map item and click **Actions > Delete**.
- If you are finished creating or editing route-map items, click **Save**.

You are returned to the **Route-maps** page.

Set rules

1. Choose **Set** as the **Rule type**.
2. In the **Rule** area, choose the rule for your set rule.

This table provides the information necessary to configure the different set rules.

Set rule type	Description	Options
set-community	Used to modify the BGP community attributes of routes. This action allows you to tag routes with specific BGP community values during routing operations, which can then be used for filtering, prioritization, traffic engineering, or other routing policies.	In the Community Numbers in ASN2:NN Format field, enter the appropriate set-community values in the ASN:NN format.
set-extcomm-list	Used to modify the BGP extended community attributes of routes by removing specific extended communities. It works by referencing a predefined extended community list and applying the action to the selected routes that match the conditions of the route map.	In the ExtCommunity Name to Delete field, choose the extended community-list to delete or click + Create extended community-list to create a new one. See Create a community-list for more information.

Set rule type	Description	Options
set-ip-next-hop set-ipv6-next-hop	Used to modify the next-hop IP address of a route. This action enables you to redirect traffic by specifying a new next-hop IP address for packets matching certain conditions defined in the route map. It is a critical feature for implementing advanced routing policies, such as policy-based routing (PBR), traffic engineering, or route redistribution.	In the Parameters area, make these configurations: 1. Next Hop Addresses: Enter the next hop addresses to be used in this set rule. You can enter multiple comma-separated next-hops in this field. 2. Put a check in the box for one or more of these options, if necessary: ◦ Drop on Failure ◦ Enable Load Sharing ◦ Force Order ◦ Verify Availability ◦ Use Peer Address ◦ Redistribute Unchanged ◦ Unchanged 3. In the Track ID field, enter a track ID, if necessary. See Configure a track for more information.
set-local-preference	Used to modify the BGP local preference (LocalPref) attribute of routes. The local preference is a well-known BGP path selection attribute used to influence outbound traffic by determining which path is preferred for sending traffic out of an Autonomous System (AS). The higher the local preference value, the more preferred the route.	In the Preference Value field, enter the appropriate preference value.

3. Click **Save**.

You are returned to the **Add route-map item** page.

4. Make additional configurations in the **Add route-map item** page, if necessary.

- If you want to create another route-map entry, click **+ Add route-map entry**.
- If you want to delete the route-map entry that you've already created, click **Delete** in that route-map entry area.
- If you are finished adding the route-map list item, click **Save**.

You are returned to the **Create route-map** page, with the newly-created route-map list item

listed.

5. Make additional configurations in the **Create route-map** page, if necessary.

- If you want to create another route-map item, click **Actions > Add route-map item**.
- If you want to edit an existing route-map item, choose the route-map item and click **Actions > Edit**.
- If you want to delete an existing route-map item, choose the route-map item and click **Actions > Delete**.
- If you are finished creating or editing route-map items, click **Save**.

You are returned to the **Route-maps** page.

Edit a route-map

To edit a route-map:

1. [Navigate to the Routing policies page](#).
2. Click the **Route-maps** subtab.
3. Choose a route-map from the configured route-maps listed in the **Route-maps** page.
4. Choose **Actions > Edit**.
5. Make the necessary changes in the appropriate pages in the configured route-map, then click **Save**.

Delete a route-map

To delete a route-map:

1. [Navigate to the Routing policies page](#).
2. Click the **Route-maps** subtab.
3. Choose a route-map from the configured route-maps listed in the **Route-maps** page.
4. Choose **Actions > Delete**.

ACLs

An Access Control List (ACL) is an ordered set of rules that you can use to filter traffic. Each rule specifies a set of conditions that a packet must satisfy to match the rule.

To view information about ACLs in your Nexus Dashboard:

1. [Navigate to the Routing policies page](#).
2. Click the **ACLs** subtab.

Any already-configured ACLs are displayed with this information:

Field	Description
Name	Shows the name of the ACL. Click the link in the Name column to display information for this ACL.
Type	Shows the type of ACL.
Entries	Shows the number of entries in this ACL.
Associations	Shows the number of associations in this ACL.

3. Configure an ACL, if necessary.

These options are available to configure ACLs:

- [Create an ACL](#)
- [Edit an ACL](#)
- [Delete an ACL](#)

Create an ACL

To create a route-map:

1. [Navigate to the Routing policies page.](#)
2. Click the **ACLs** subtab.
3. Choose **Actions > Create ACL**.
4. Enter a name for the access control list in the **Name** field.
5. In the **Type** field, choose **IPv4** or **IPv6**.
6. Review the existing access-list entries, if any are already configured.

Any already-configured access-list entries are displayed with this information:

Field	Description
Sequence number	Shows the sequence to insert to or delete from the existing access-list.
Action	Shows the action that this access-list entry will take: ▪ permit ▪ deny ▪ remark

Field	Description
IP protocol/options	Shows the IP protocol/option: ▪ TCP ▪ ICMP ▪ IP ▪ UDP ▪ EIGRP ▪ OSPF ▪ PIM ▪ IGMP ▪ Custom: Provide the value for the Custom Protocol Number when choosing this option. ○ You must provide a protocol number between 0-255. ○ You cannot use the following reserved protocol numbers for the Custom Protocol Number: ▪ IPv4: 0, 1, 2, 6, 17, 47, 50, 51, 88, 89, 94, 103, 108 ▪ IPv6: 1, 6, 17, 41, 50, 51, 88, 103, 108
Source port match operator	Shows the source port match operator: ▪ equalTo ▪ greaterThan ▪ lessThan ▪ notEqualTo ▪ portRange ▪ none
Source port	Shows the source port.
Source port range start	Shows the port range start.
Source port range end	Shows the source port range end.
Source address	Shows the source IP address.
Destination port match operator	Shows the destination port match operator: ▪ equalTo ▪ greaterThan ▪ lessThan ▪ notEqualTo ▪ portRange ▪ none
Destination port	Shows the destination port.

Field	Description
Destination port range start	Shows the destination port range start.
Destination port range end	Shows the destination port range end.
Destination address	Shows the destination IP address.
TCP session rules	Shows the TCP session rule: ▪ ack ▪ fin ▪ established ▪ psh ▪ rst ▪ syn
ICMP options	Shows the ICMP option: ▪ Administratively Prohibited ▪ Echo ▪ Echo Reply ▪ General Parameter Problem ▪ Host Isolated ▪ Host Precedence Unreachable ▪ Host Redirect ▪ Host Tos Unreachable

7. In the **Access-list entries** area, click **Actions > Create ACL entry**.

8. In the **Create ACL entry** page, enter the following information:

Field	Description
Sequence number	Enter the sequence to insert to the existing access-list entry. Valid value is between 1-4294967295.
Action	Enter the action that this access-list entry will take: ▪ permit ▪ deny ▪ remark

Field	Description
IP protocol/options	Choose the IP protocol/option: ▪ TCP ▪ ICMP ▪ IP ▪ UDP ▪ EIGRP ▪ OSPF ▪ PIM ▪ IGMP ▪ Custom: Provide the value for the Custom Protocol Number when choosing this option. ○ You must provide a protocol number between 0-255. ○ You cannot use the following reserved protocol numbers for the Custom Protocol Number: ▪ IPv4: 0, 1, 2, 6, 17, 47, 50, 51, 88, 89, 94, 103, 108 ▪ IPv6: 1, 6, 17, 41, 50, 51, 88, 103, 108
Source port match operator	Displayed if you chose TCP or UDP as the IP protocol/option. Choose the source port match operator: ▪ equalTo ▪ greaterThan ▪ lessThan ▪ notEqualTo ▪ portRange ▪ none
Source port	Displayed if you chose TCP or UDP as the IP protocol/option and any option other than portRange as the Source port match operator. Enter the source port. It must be a valid port number (0-65535) or service name. Examples of valid entries include: bgp chargen cmd daytime discard domain drip echo exec finger ftp ftp-data gopher hostname ident irc klogin kshell login lpd nnntp pim-auto-rp pop2 pop3 smtp sunrpc tacacs talk telnet time uucp www ([0-9] [1-9][0-9]{1,3} [1-5][0-9]{4} 6[0-4][0-9]{3} 65[0-4][0-9]{2} 655[0-2][0-9] 6553[0-5])

Field	Description
Source port range start	Displayed if you chose TCP or UDP as the IP protocol/option and portRange as the Source port match operator. Enter the source port range start. It must be a valid port number (0-65535) or service name. Examples of valid entries include: bgp chargen cmd daytime discard domain drip echo exec finger ftp ftp-data gopher hostname ident irc klogin kshell login lpd nnntp pim-auto-rp pop2 pop3 smtp sunrpc tacacs talk telnet time uucp www ([0-9] [1-9][0-9]{1,3} [1-5][0-9]{4} 6[0-4][0-9]{3} 65[0-4][0-9]{2} 655[0-2][0-9] 6553[0-5])
Source port range end	Displayed if you chose TCP or UDP as the IP protocol/option and portRange as the Source port match operator. Enter the source port range end. It must be a valid port number (0-65535) or service name. Examples of valid entries include: bgp chargen cmd daytime discard domain drip echo exec finger ftp ftp-data gopher hostname ident irc klogin kshell login lpd nnntp pim-auto-rp pop2 pop3 smtp sunrpc tacacs talk telnet time uucp www ([0-9] [1-9][0-9]{1,3} [1-5][0-9]{4} 6[0-4][0-9]{3} 65[0-4][0-9]{2} 655[0-2][0-9] 6553[0-5])
Source address	Displayed if you chose any option other than TCP or UDP as the IP protocol/option. Enter the source IP address. Source IP address must be either: • any • A valid IPv4 address, preferably with CIDR (for example, 192.168.1.1, 10.0.0.0/8) • A valid IPv6 address, preferably with CIDR (for example, 2001:db8::1, 2001:db8::/32)
Destination port match operator	Displayed if you chose TCP or UDP as the IP protocol/option. Choose the destination port match operator: • equalTo • greaterThan • lessThan • notEqualTo • portRange • none

Field	Description
Destination port	Displayed if you chose TCP or UDP as the IP protocol/option and any option other than portRange as the Source port match operator. Enter the destination port. It must be a valid port number (0-65535) or service name. Examples of valid entries include: bgp chargen cmd daytime discard domain drip echo exec finger ftp ftp-data gopher hostname ident irc klogin kshell login lpd nnntp pim-auto-rp pop2 pop3 smtp sunrpc tacacs talk telnet time uucp www ([0-9] ([1-9][0-9]{1,3} [1-5][0-9]{4} 6[0-4][0-9]{3} 65[0-4][0-9]{2} 655[0-2][0-9] 6553[0-5]))
Destination port range start	Displayed if you chose TCP or UDP as the IP protocol/option and any option other than portRange as the Destination port match operator. Enter the destination port range start. It must be a valid port number (0-65535) or service name. Examples of valid entries include: bgp chargen cmd daytime discard domain drip echo exec finger ftp ftp-data gopher hostname ident irc klogin kshell login lpd nnntp pim-auto-rp pop2 pop3 smtp sunrpc tacacs talk telnet time uucp www ([0-9] ([1-9][0-9]{1,3} [1-5][0-9]{4} 6[0-4][0-9]{3} 65[0-4][0-9]{2} 655[0-2][0-9] 6553[0-5]))
Destination port range end	Displayed if you chose TCP or UDP as the IP protocol/option and any option other than portRange as the Destination port match operator. Enter the destination port range end. It must be a valid port number (0-65535) or service name. Examples of valid entries include: bgp chargen cmd daytime discard domain drip echo exec finger ftp ftp-data gopher hostname ident irc klogin kshell login lpd nnntp pim-auto-rp pop2 pop3 smtp sunrpc tacacs talk telnet time uucp www ([0-9] ([1-9][0-9]{1,3} [1-5][0-9]{4} 6[0-4][0-9]{3} 65[0-4][0-9]{2} 655[0-2][0-9] 6553[0-5]))
Destination address	Displayed if you chose any option other than TCP or UDP as the IP protocol/option. Enter the destination IP address. Destination IP address must be either: • any • A valid IPv4 address, preferably with CIDR (for example, 192.168.1.1, 10.0.0.0/8) • A valid IPv6 address, preferably with CIDR (for example, 2001:db8::1, 2001:db8::/32)

Field	Description
TCP flags	Displayed if you chose TCP as the IP protocol/option. Choose the TCP session rule: ▪ ack ▪ fin ▪ established ▪ psh ▪ rst ▪ syn
ICMP options	Displayed if you chose ICMP as the IP protocol/option. Choose the ICMP option. ▪ If you chose IPv4 in the Type field, these options are available: Beyond Scope, Destination Unreachable, Echo Request, Echo Reply, Hop Limit, Mid Query, Mid Report, Mid Reduction, Mldv2, Nd Ns, Nd Na, No Admin, No Route, Packet Too Big, Parameter Problem, Port Unreachable, Reassembly Timeout, Router Advertisement, Router Solicitation, Unreachable ▪ If you chose IPv6 in the Type field, these options are available: Administratively Prohibited, Echo, Echo Reply, General Parameter Problem, Host Isolated, Host Precedence Unreachable, Host Redirect, Host Tos Unreachable, Host Unknown, Host Unreachable, Net Unreachable, Network Unknown, Option Missing, Packet Too Big, Parameter Problem, Port Unreachable, Router Advertisement, Router Solicitation, TTL Exceeded, Unreachable



Nexus Dashboard will perform checks on certain options and will not allow you to continue if the necessary information is not provided in those options.

9. Click **Save**.

You are returned to the **Create ACL** page.

10. Verify that the information displayed for the newly-created access control list is correct, then click **Save**.

You are returned to the **ACLs** page.

Edit an ACL

To edit an ACL:

1. [Navigate to the Routing policies page](#).

2. Click the **ACLs** subtab.
3. Choose an ACL from the configured ACLs listed in the **ACLs** page.
4. Choose **Actions > Edit**.
5. Make the necessary changes in the appropriate pages in the configured ACL, then click **Save**.

Delete an ACL

To delete an ACL:

1. [Navigate to the Routing policies page](#).
2. Click the **ACLs** subtab.
3. Choose an ACL from the configured ACLs listed in the **ACLs** page.
4. Choose **Actions > Delete**.

Prefix-list

A prefix list is a mechanism to define and manage IP address ranges or prefixes. These lists are primarily employed for configuring network policies, such as route filtering, access control, or traffic redirection in the Nexus Dashboard fabric.

To view information about prefix-lists in your Nexus Dashboard:

1. [Navigate to the Routing policies page](#).
2. Click the **Prefix-list** subtab.

Any already-configured prefix-lists are displayed with this information:

Field	Description
Name	Shows the name of the prefix-list. Click the link in the Name column to display information for this prefix-list.
Type	Shows the type of prefix-list.
Entries	Shows the number of entries in this prefix-list.
Associations	Shows the number of associations in this prefix-list.

3. Configure a prefix-list, if necessary.

These options are available to configure prefix-lists:

- [Create a prefix-list](#)
- [Edit a prefix-list](#)
- [Delete a prefix-list](#)

Create a prefix-list

To create a prefix-list:

1. [Navigate to the Routing policies page.](#)
2. Click the **Prefix-list** subtab.
3. Choose **Actions > Create prefix-list**.
4. Enter a name for the prefix-list in the **Name** field.
5. In the **Type** field, choose **IPv4** or **IPv6**.
6. In the **Description** field, enter a description for the prefix-list, if necessary.
7. In the **Prefix-list entries** area, click **Actions > Create prefix-list entry**.
8. In the **Create prefix-list entry** page, enter the following information:

Field	Description
Sequence number	Enter the sequence to insert to the existing prefix-list entry. Valid value is between 1-4294967294.
Action	Enter the action that this prefix-list entry will take: ▪ permit ▪ deny
Prefix	Enter the IP prefix network/length (for example, 192.168.1.0/24).
Exact prefix length	Enter the exact prefix length to be matched.
Minimum prefix length	Enter the minimum prefix length to be matched.
Maximum prefix length	Enter the maximum prefix length to be matched.
Explicit match mask	Enter the mask that specifies the bits of a prefix address in a prefix-list that are compared to, in format A.B.C.D .

9. Click **Save**.

You are returned to the **Create prefix-list** page.

10. Verify that the information displayed for the newly-created prefix-list entry is correct, then click **Save**.

You are returned to the **Prefix-list** page.

Edit a prefix-list

To edit a prefix-list:

1. [Navigate to the Routing policies page.](#)
2. Click the **Prefix-list** subtab.
3. Choose a prefix-list from the configured route-maps listed in the **Prefix-list** page.
4. Choose **Actions > Edit**.
5. Make the necessary changes in the appropriate pages in the configured prefix-list, then click **Save**.

Delete a prefix-list

To delete a prefix-list:

1. [Navigate to the Routing policies page](#).
2. Click the **Prefix-list** subtab.
3. Choose a prefix-list from the configured route-maps listed in the **Route-maps** page.
4. Choose **Actions > Delete**.
5. In the **Warning** message that pops up, click **Confirm**.

About community lists and extended community lists

Under **Routing policies** in Nexus Dashboard, there are two types of community-lists available:

- Community-lists
- Extended community-lists

The primary difference between **community lists** and **extended community lists** lies in the type of attributes they match and how they are used in BGP (Border Gateway Protocol) routing policies.

Here's a detailed comparison:

Key differences between community-lists and extended community lists

Feature	Community lists	Expanded community lists
Purpose	Match standard BGP community attributes, which are 32-bit values used to tag routes for filtering, redistribution, or prioritization.	Match extended BGP community attributes, which are 64-bit values providing additional tagging capabilities for more advanced use cases, such as VPN route targets, QoS, or Site-of-Origin (SOO).
Attribute size	Operate on 32-bit community attributes. Examples: 65000:100 or 100:1 .	Operate on 64-bit extended community attributes. Examples: RT:65000:100 (Route Target) or SOO:65000:200 (Site-of-Origin).
Use cases	Used for general route tagging and filtering based on standard BGP community values. Examples: Route prioritization, route filtering, or marking routes to influence path selection.	Used for more complex scenarios requiring additional tagging information. Examples: MPLS VPNs (to match Route Targets), Site-of-Origin tagging, QoS policies, traffic engineering.

Types	- Standard community lists: Match exact 32-bit community values. - Expanded community lists: Match patterns of 32-bit community values using regex.	- Standard extended community lists: Match exact 64-bit extended community values. - Expanded extended community lists: Match patterns of 64-bit extended community attributes using regex.
Flexibility	Limited to matching standard 32-bit communities, which are less flexible for complex scenarios.	Provide greater flexibility due to their ability to match 64-bit extended communities and support advanced use cases like MPLS VPNs or SOO tagging.

Refer to these sections for configuration procedures for each of these types of community-lists:

- [Community-list](#)
- [Extended community-list](#)

Community-list

Community-lists are used to define and group Border Gateway Protocol (BGP) community attributes. These attributes allow you to tag and group routes in a BGP environment, enabling more granular control over how routes are distributed, filtered, or prioritized within the Nexus Dashboard fabric or between external networks.

In Nexus Dashboard, **standard community lists** and **expanded community lists** are two types of BGP community lists used to manage and filter routes based on their BGP community attributes. The primary difference lies in the way they match community attributes and their flexibility.

Key differences between standard and expanded community lists

Feature	Standard Community Lists	Expanded Community Lists
Matching method	Matches specific BGP community values exactly.	Matches BGP community values using regular expressions (regex).
Flexibility	Less flexible; only exact matches are supported.	More flexible; can match patterns or ranges of community values.
Examples of usage	Useful for simple, predefined community tags.	Useful for more complex scenarios where multiple patterns need to be matched.
Configuration simplicity	Easier to configure, as it only requires the exact value.	More complex configuration due to the use of regex.

To view information about standard or expanded community-lists in your Nexus Dashboard:

1. [Navigate to the Routing policies page](#).
2. Click the **Community-list** subtab.

Any already-configured community-lists are displayed with this information:

Field	Description
Name	Shows the name of the community-list. Click the link in the Name column to display information for this community-list.
Type	Shows the type of community-list.
Entries	Shows the number of entries in this community-list.
Associations	Shows the number of associations in this community-list.

3. Configure a community-list, if necessary.

These options are available to configure community-list:

- [Create a community-list](#)
- [Edit a community-list](#)
- [Delete a community-list](#)

Create a community-list

To create a community-list:

1. [Navigate to the Routing policies page.](#)
2. Click the **Community-list** subtab.
3. Choose **Actions > Create community-list**.
4. Enter a name for the community-list in the **Name** field.
5. In the **Type** field, choose **standard** or **expanded**.

The information displayed in the **Create community-list** page, and the create action, varies depending on the type of community-list that you choose.

- [Create standard community-list entry](#)
- [Create expanded community-list entry](#)

Create standard community-list entry

1. Review the existing standard community-list entries, if any are already configured.

Any already-configured standard community-list entries are displayed with this information:

Field	Description
Sequence number	Shows the sequence to insert to the existing community-list entry.
Action	Shows the action that this community-list entry will take: ▪ permit ▪ deny
No-advertise	Shows if the No-advertise option is enabled.
Blackhole	Shows if the Blackhole option is enabled.

Field	Description
No-export	Shows if the No-export option is enabled.
Internet	Shows if the Internet option is enabled.
Graceful shutdown	Shows if the Graceful shutdown option is enabled.
Local-As	Shows if the Local-As option is enabled.
Community numbers	Shows the community numbers, in the ASN:NN format.

- In the **Standard community-list entries** area, click **Actions > Create standard community-list entry**.
- In the **Create standard community-list entry** page, enter the following information:

Field	Description
Sequence number	Enter the sequence to insert to the existing community-list entry. Valid value is between 1-4294967294.
Action	Enter the action that this community-list entry will take: • permit • deny
No-advertise	Put a check in the box to not advertise to any peer.
Blackhole	Put a check in the box to enable the blackhole option. BGP blackhole routing is a destination-based packet filtering technology that uses the Border Gateway Protocol (BGP) to “advertise” or communicate a blackhole route for a specific IP address to other routers within the BGP network.  The blackhole option is supported on switches running NX-OS 10.3(2) or later.
No-export	Put a check in the box to not export to next AS.
Internet	Put a check in the box to enable the internet option.
Graceful shutdown	Put a check in the box to signal the graceful shutdown of paths. Use this community to reduce the amount of traffic lost when BGP peering sessions are about to be shut down deliberately (for example, for planned maintenance).
Local-As	Put a check in the box to not send outside local AS.
Community numbers	Enter the community numbers, in the ASN:NN format.

- Click **Save**.

You are returned to the **Create community-list** page.

- Verify that the information displayed for the newly-created community-list entry is correct, then click **Save**.

You are returned to the **Community-list** page.

Create expanded community-list entry

1. Review the existing expanded community-list entries, if any are already configured.

Any already-configured expanded community-list entries are displayed with this information:

Field	Description
Sequence number	Shows the sequence to insert to the existing expanded community-list entry.
Action	Shows the action that this expanded community-list entry will take: ▪ permit ▪ deny
Community number REGEX	Shows the community number REGEX value.

2. In the **Expanded community-list entries** area, click **Actions > Create expanded community-list entry**.
3. In the **Create expanded community-list entry** page, enter the following information:

Field	Description
Sequence number	Enter the sequence to insert to the existing expanded community-list entry. Valid value is between 1-4294967294.
Action	Enter the action that this expanded community-list entry will take: ▪ permit ▪ deny
Community number REGEX	Enter the community number regular-expression (REGEX) in aa:nn format.

4. Click **Save**.

You are returned to the **Create community-list** page.

5. Verify that the information displayed for the newly-created community-list entry is correct, then click **Save**.

You are returned to the **Community-list** page.

Edit a community-list

To edit a community-list:

1. [Navigate to the Routing policies page](#).
2. Click the **Community-list** subtab.
3. Choose a community-list from the configured community-lists shown in the **Community-list**

page.

4. Choose **Actions > Edit**.
5. Make the necessary changes in the appropriate pages in the configured community-list, then click **Save**.

Delete a community-list

To delete a community-list:

1. [Navigate to the Routing policies page](#).
2. Click the **Community-list** subtab.
3. Choose a community-list from the configured community-lists shown in the **Community-list** page.
4. Choose **Actions > Delete**.
5. In the **Warning** message that pops up, click **Confirm**.

Extended community-list

Extended community lists are used to match BGP extended community attributes, which provide additional tagging and control over BGP routes. Extended communities are an enhancement to the standard BGP communities, and offer more flexibility and granularity in applying routing policies.

Extended communities are 64-bit values (compared to the 32-bit community attributes) that allow for more complex route tagging. They are commonly used in scenarios such as:

- VPN route target (RT) attributes in MPLS Layer 3 VPNs.
- Route origin information.
- Site-of-origin (SOO) tagging to prevent routing loops.
- Traffic engineering or QoS policies.

Similar to community lists as described in [Community-list](#), extended community lists can be of two types:

- Standard extended community lists:
 - Matches specific extended community values exactly.
 - Best suited for scenarios where you know the exact extended community attributes.
- Expanded extended community lists:
 - Uses regular expressions (regex) to match patterns of extended community attributes.
 - Useful when dealing with ranges or subsets of extended communities.

To view information about extended community-lists in your Nexus Dashboard:

1. [Navigate to the Routing policies page](#).
2. Click the **Extended community-list** subtab.

Any already-configured extended community-lists are displayed with this information:

Field	Description
Name	Shows the name of the extended community-list. Click the link in the Name column to display information for this community-list.
Type	Shows the type of extended community-list.
Entries	Shows the number of entries in this extended community-list.
Associations	Shows the number of associations in this extended community-list.

3. Configure an extended community-list, if necessary.

These options are available to configure extended community-lists:

- [Create an extended community-list](#)
- [Edit an extended community-list](#)
- [Delete an extended community-list](#)

Create an extended community-list

To create an extended community-list:

1. [Navigate to the Routing policies page.](#)
2. Click the **Extended community-list** subtab.
3. Choose **Actions > Create extended community-list**.
4. Enter a name for the extended community-list in the **Name** field.
5. In the **Type** field, choose **standard** or **expanded**.

The information displayed in the **Create community-list** page, and the create action, varies depending on the type of community-list that you choose.

- [Create standard extended community-list entry](#)
- [Create expanded extended community-list entry](#)

Create standard extended community-list entry

1. Review the existing standard extended community-list entries, if any are already configured.

Any already-configured standard extended community-list entries are displayed with this information:

Field	Description
Sequence number	Shows the sequence to insert to the existing standard extended community-list entry.
Action	Shows the action that this standard extended community-list entry will take: ▪ permit ▪ deny

Field	Description
Router MAC list	Shows the router MAC list.
Route target list	Shows the route target list.
Site-of-Origin list	Shows the site-of-origin list.
Transitive generic extended community	Shows the transitive generic extended community.
Non-transitive generic extended community	Shows the non-transitive generic extended community.

- In the **Standard extended community-list entries** area, click **Actions > Create standard extended community-list entry**.
- In the **Create standard extended community-list entry** page, enter the following information:

Field	Description
Sequence number	Enter the sequence to insert to the existing standard extended community-list entry. Valid value is between 1-4294967294.
Action	Enter the action that this standard extended community-list entry will take: • permit • deny
Router MAC list	Enter a comma-separated router MAC list. For example: d478:111.57b8, d478.1111.2222
Route target list	Enter a comma-separated route target list. For example: 1:2, 2:2
Site-of-Origin list	Enter a comma-separated site-of-origin list. For example: 11:12, 15:15
Transitive generic extended community	Enter the transitive generic extended community. For example: 12:22
Non-transitive generic extended community	Enter the non-transitive generic extended community. For example: 10:10, 11:11

- Click **Save**.

You are returned to the **Create extended community-list** page.

- Verify that the information displayed for the newly-created extended community-list entry is correct, then click **Save**.

You are returned to the **Extended community-list** page.

Create expanded extended community-list entry

1. Review the existing expanded extended community-list entries, if any are already configured.

Any already-configured expanded extended community-list entries are displayed with this information:

Field	Description
Sequence number	Shows the sequence to insert to the existing expanded extended community-list entry.
Action	Shows the action that this expanded extended community-list entry will take: ▪ permit ▪ deny
Community number REGEX	Shows the community number REGEX value.

2. In the **Expanded extended community-list entries** area, click **Actions > Create expanded extended community-list entry**.
3. In the **Create expanded extended community-list entry** page, enter the following information:

Field	Description
Sequence number	Enter the sequence to insert to the existing expanded extended community-list entry. Valid value is between 1-4294967294.
Action	Enter the action that this expanded extended community-list entry will take: ▪ permit ▪ deny
Community number REGEX	Enter the community number regular-expression (REGEX) in aa:nn format.

4. Click **Save**.

You are returned to the **Create extended community-list** page.

5. Verify that the information displayed for the newly-created extended community-list entry is correct, then click **Save**.

You are returned to the **Extended community-list** page.

Edit an extended community-list

To edit an extended community-list:

1. [Navigate to the Routing policies page](#).
2. Click the **Extended community-list** subtab.

3. Choose an extended community-list from the configured extended community-lists listed in the **Extended community-list** page.
4. Choose **Actions > Edit**.
5. Make the necessary changes in the appropriate pages in the configured extended community-list, then click **Save**.

Delete an extended community-list

To delete an extended community-list:

1. [Navigate to the Routing policies page](#).
2. Click the **Extended community-list** subtab.
3. Choose an extended community-list from the configured extended community-lists listed in the **Extended community-list** page.
4. Choose **Actions > Delete**.
5. In the **Warning** message that pops up, click **Confirm**.

Inter-Fabric

To work with inter-fabric connectivity in a fabric:

1. [Navigate to the Connectivity page](#).
2. Click the **Inter-Fabric** tab.

A page showing the information for previously-configured inter-fabric connections appears.

The following table describes the fields that appear on the **Inter-Fabric** tab.

Field	Description
Status	Provides one of the following states as the status of the inter-fabric connections: ▪ Connected ▪ Degraded ▪ Failed
Physical Links	Shows how many physical links in the inter-fabric connections are up or down.
Overlay BGP Peers	Shows how many overlay BGP peers in the inter-fabric connections are up or down.
Overlay Tunnels	Shows how many overlay tunnels in the inter-fabric connections are up or down.
Fabric Type	Shows the fabric types used in the inter-fabric connections.

L3 Neighbors

Nexus Dashboard provides a unified view of all L3Outs and their neighbors in your fabric. This information provides visibility into operational data reported by the fabrics controller about fabric-level connectivity and simplifies troubleshooting by showing the various Layer 3 adjacencies (neighbors) for each L3Out.

Use the **L3 Neighbors** page to view L3 neighbors for this fabric. You can filter the results based on neighbor, local switch, routing protocol, VRF, and operational status. OSPF and OSPFv3 support includes OSPF and OSPFv3 statistics, operational statistics, interface statistics, and neighbor statistics. Click the IP address in the **Neighbor** column to view details on this neighbor.

Nexus Dashboard updates information on this page in real-time upon modification or change of the properties.

View L3 neighbors in a fabric

Follow these steps to view L3 neighbors in a fabric.

1. [Navigate to the Connectivity page.](#)
2. Click the **L3 Neighbors** tab.

A page showing the information for previously-configured L3 neighbors appears.

The **L3 Neighbors** page provides a unified view of all neighbors based on L3Out configuration for that fabric in tabular format. You can **Filter** or sort the table based on each column.

3. Click an entry in the **Neighbor** column to view that neighbor's details.

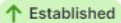
Here, you can view the **Local Switch** information (including its name, IP address, ASN, interface information, and so forth) and **Neighbor Details** (such as its IP address, ASN, route ID, port, and so forth).

For example, the following figure shows sample information for BGP L3Out neighbors:

Local Switch Details

Name	Local IP	ASN	Interface	Router ID	Port	VRF
NDI-N06-LEAF2	10.2.0.2	3000	lo0	10.2.0.2	21974	default

NEIGHBOR DETAILS

Neighbor IP	ASN	Router ID	Port	Neighbor Status	Uptime
10.2.0.5	3000	10.2.0.5	179	 Established	7 weeks, 5 hours, 36 minutes

BGP Address Families

Neighbor Capabilities

Capability	Advertised	Received
as4	 Yes	 Yes
cap	 No	 Yes
dynamic	 Yes	 Yes

4. If the displayed information is not accurate, verify L3Out configurations.

If the L3Out neighbors are not present in the table view:

- Verify that the L3Out policy is configured in Nexus Dashboard and deployed successfully. The information is displayed only for L3Outs that are configured in Nexus Dashboard.
- Verify that the L3Out neighbors are present in Nexus Dashboard's inventory using the APIs.
 - For BGP: `GET /mso/api/v1/inventorybgpneighbors?status.fabric=<fabric-id>`
 - For OSPF: `GET /mso/api/v1/inventoryospfneighbors?status.fabric=<fabric-id>`

If the L3Out neighbors' operational state is not green:

- Verify that the switch interfaces are not in the **shut** state on either of the switches.
- Verify that the protocol settings are configured correctly and there is no mismatch in the peer device configuration.
 - For BGP, check the authentication, eBGP MultiHop TTL, and ASN are configured correctly.
 - For OSPF, check authentication, Area ID, and MTU configurations.

Guidelines and limitations for L3 neighbors

These guidelines and limitations apply for L3 neighbors:

- In an NX-OS fabric, OSPFv3 is supported on switches with the Cisco NX-OS 10.5.x release and later.
- Up to 20 GUI sessions at one time can see real-time updates. You must close one of the previous sessions for a new session to be able to see real-time updates.

Endpoints

The **Endpoints** page provides information on the endpoints for this fabric. You can filter the results based on MAC address, IP addresses, hostname, anomaly level, interface, status, VRF instance, Layer 3 switches, and Layer 3 interfaces.

Nexus Dashboard updates the information on this page in real-time upon modification or change of the properties.

View endpoints in a fabric

Follow these steps to view endpoints in a fabric.

1. [Navigate to the Connectivity page](#).
2. Click the **Endpoints** tab.

A page showing the information for previously-configured endpoints appears.

3. Click a MAC address in the **MAC** column to get the following additional information on that endpoint:
 - Overview
 - General
 - VM Name
 - Hypervisor
 - MAC Address
 - IP Address
 - Hostname
 - Last Updated
 - Status
 - Network Configuration
 - Tenant
 - VRF
 - EPG/L3 Out
 - BD
 - Encap
 - Connected To
 - Nodes
 - Interface
 - Endpoint history—Determine how you want to show endpoint history. The time range for the history is only based on the time selector chosen at the fabric level.
 - Anomalies

You can also click the IP address or the hostname (if clickable) to view additional information about them.

Guidelines and limitations for endpoints

- Endpoint page updates can take up to 60 seconds and therefore are not always in real-time.
- A maximum of 20 GUI sessions can view real-time updates simultaneously. To enable a new session to access real-time updates, you must close one of the previous sessions.

Routes

The **Routes** page provides information on the routes for this fabric. You can filter the results based on route, protocol, VRF instance, or next hop.

Nexus Dashboard updates information on this page in real-time upon modification or change of the properties.

View routes in a fabric

Follow these steps to view routes in a fabric.

1. [Navigate to the Connectivity page](#).
2. Click the **Routes** tab.

A page showing the information for previously-configured routes appears.

3. Choose whether you want to view the details for IPv4, IPv6, or multicast routes.

View IPv4 or IPv6 routes

Click the **IPv4** or **IPv6** vertical tab.

- Click **Download** to download any route table from the last 7 days.
- You can filter the table based on route, protocol or VRF.
- The **Routes per VRF** graph shows the number of routes per VRF for the top 10 VRFs.
- The **Change Summary** helps view the number of the added, modified, and deleted routes. Click **Change Summary** to list all the routes which have had modified or deleted change made in a tabular form.
- The Routes table displays information such as route, node name, protocol, VRF, and Next Hop. The Routes table only displays 10,000 route table entries and 10,000 history events. The Routes table always shows the current route table irrespective of the time range selected.
- Click on any route listed to view General and Route History view along with the timeline for the change. The General page list of all the next hops available for that route. The Route History page shows detail of last 50 route events. The route history events shows deleted or modified events. Route add events are shown under modified routes category.
- Each Next Hop IP address is subscripted with the node name and interface where it is configured.
- Click **View Details** in the routing page banner to list any unsupported nodes, inactive nodes, switch connection, failure and if data is inconsistent for the selected time range.

Guidelines and limitations for routes

These guidelines and limitations apply for routes:

- Routes information for NX-OS switches is updated in real-time.
- A maximum of 20 GUI sessions can view real-time updates simultaneously. To enable a new session to access real-time updates, you must close one of the previous sessions.

- Nexus Dashboard does not have support for ACI routes in Nexus Dashboard release 4.1.1.

Multicast traffic monitoring overview

Nexus Dashboard provides the following two views for monitoring multicast traffic in your fabric.

- Multicast routes—provides multicast routing state (control plane), including active (S,G) and (*,G) entries, outgoing interface lists, and detailed source and receiver attachments. You can use this view to verify multicast configuration, validate tree-building behavior, and visualize the end-to-end flow path from the first-hop router (FHR) to the last-hop router (LHR).
- Multicast flows—provides live traffic telemetry (data plane) for active multicast traffic, including multicast group, tenant and VRF context, source attachment details, and source traffic rate. You can use this view to monitor live traffic at the flow-table level alongside unicast flow data.

Both views draw from the same multicast telemetry data exported by Cisco N9000 Series switches. Beginning with release 4.3.1, Nexus Dashboard also provides end-to-end flow path topology visualization from the FHR to the LHR. For more information, see [Multicast telemetry flow path visibility](#).

Multicast routes

Multicast telemetry flow path visibility

Beginning with release 4.3.1, Nexus Dashboard introduces multicast telemetry flow path visibility feature that provides end-to-end traceability for multicast traffic within your fabric. Nexus Dashboard exports multicast routes and state information from Cisco N9000 Series switches, enabling you to visualize the complete path a multicast flow takes from the source to each individual receiver.

This feature provides the following capabilities.

- View the multicast flow path from the FHR to the LHR.
- Access granular statistics for multicast sources and receivers, including packet rates and byte counts.
- Correlate multicast routes with LLDP and CDP neighbor information to identify the exact physical path, including sub-interfaces and switched virtual interfaces (SVIs).

The multicast dashboard also helps you analyze network traffic by identifying top multicast groups based on the following criteria.

- Top groups by traffic—identifies the multicast groups consuming the most bandwidth.
- Top groups by receivers— identifies the multicast groups with the highest number of active receivers.

Prerequisites

Before you enable multicast telemetry flow path visibility, ensure that your environment meets the following requirements.

- Supported platform: Cisco N9000 Series switches in standalone mode.

- Software version: Cisco NX-OS release 10.4(1)F or later for VXLAN fabrics, or release 10.3(3)F or later for Classic LAN fabrics.
- Enable the following commands on each participating Cisco Nexus switch.

```
switch(config)# hardware profile multicast flex-stats-enable  
switch(config)# multicast flow-path export
```



To apply the changes, save the configuration and reload the switch.

```
switch(config)# copy running-config startup-config  
switch(config)# reload
```

Monitor multicast routes

Once you configure the fabric, refer the following sections to monitor multicast traffic.

View flow statistics

Flow statistics provide granular details to help you analyze traffic health and identify performance issues.

Follow these steps to view detailed multicast flow metrics.

1. Navigate to the **Fabrics** page.

Go to **Manage > Fabrics**.

2. Click the fabric you configured to collect multicast data.

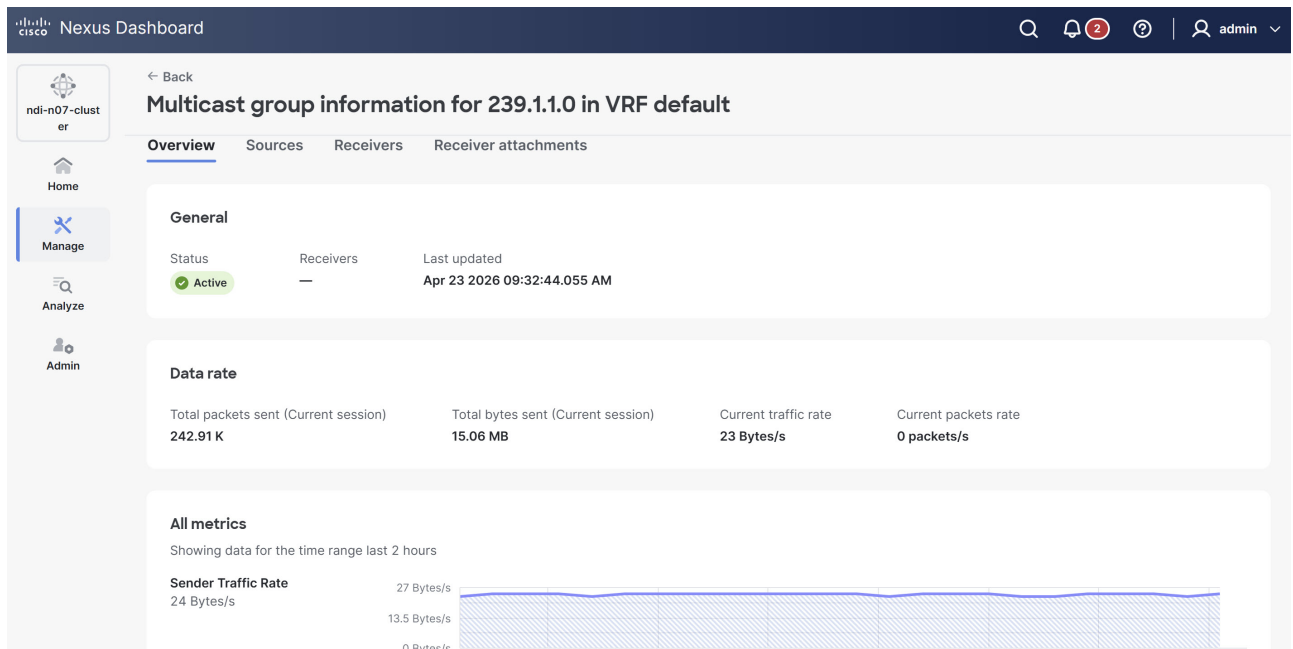
The **Fabric Overview** page displays.

3. Click the **Connectivity** tab and then click the **Routes** tab.
4. Click the **Multicast** vertical tab.

You can view the information such as multicast groups, sources sending to this group, sending rate per source, number of receivers for this group, and VRF instances.

5. Click the drop-down list to choose **Grouped** or **Ungrouped** to organize the multicast routes.
6. Click a specific multicast group to view additional information for sources and receivers.

The multicast group information page displays.



7. Click the following tabs to monitor specific multicast metrics.

- **Overview**—displays the current status (**Active/Inactive**), last update timestamp, source connection details (node and interface), session data rates, and visual trend graphs for **Sender traffic rate**, **Sender packet rate**, and **Receivers**.
- **Sources**—displays the source IP address, switch name, and interface information.
- **Receivers**— displays each receiver for the group, including receiver IP address, connected switch, interface, and join status.

View flow path topology

The flow path topology provides a visual representation of how multicast traffic traverses your network from the source to the receivers.

Follow these steps to view the flow path topology.

1. Navigate to the **Fabrics** page.

Go to **Manage > Fabrics**.

2. Click the fabric you configured to collect multicast data.

The **Fabric Overview** page displays.

3. Click the **Connectivity** tab and then click the **Routes** tab.
4. Click the **Multicast** vertical tab.

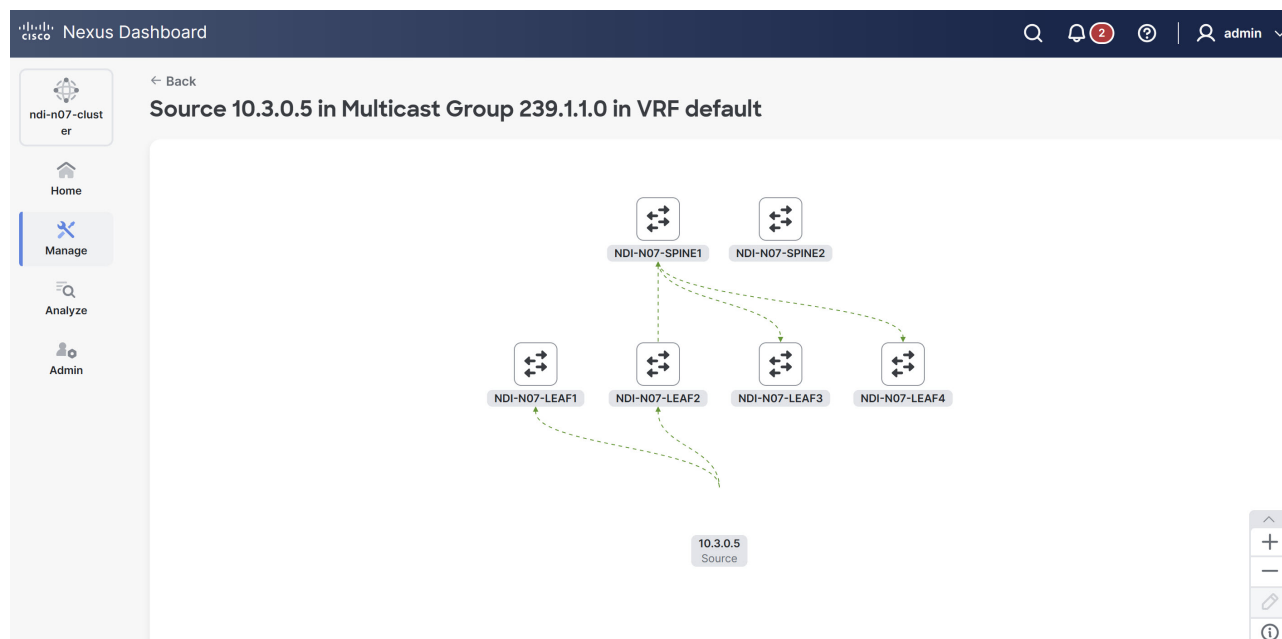
You can view the information such as multicast groups, sources sending to this group, sending rate per source, number of receivers for this group, and VRF instances.

5. Click the drop-down list to choose **Grouped** or **Ungrouped** to organize the multicast flows.
6. Click a specific multicast source.

The source associated with the multicast group page displays.

7. Click the **Receivers** tab.
8. Click **View in topology**.

Nexus Dashboard displays the end-to-end traffic flow between the source and the receivers.



Guidelines and limitations for multicast routes

The following guidelines and limitations apply to multicast routes and multicast telemetry.

Platform and protocol support

- Fabric support: Classic LAN and VXLAN fabrics are supported.
- Protocol support: Only IPv4 multicast is supported.

Topology and connectivity

- Source/receiver placement: The source and receiver must exist within the same VRF, tenant, and fabric. Both must be part of the managed fabric.
- External sources: When the FHR for multicast or the multicast source is configured outside the fabric, Nexus Dashboard does not display multicast routes.
- Connectivity support:
 - Receivers: Topology visibility is supported for ToR and access connectivity.
 - Layer 2 nodes: Topology view is not supported for sources connected to Layer 2 nodes.
- Multi-site: Multi-site topology visibility is not supported.
- VPC considerations: For VPC-connected sources, the topology view displays both nodes, and statistics appear duplicated. This occurs because the switch multicast route information does not identify the active forwarding node.

Supported scale

- Maximum groups: 2,000
- Maximum (S,G) entries: 16,000
- Maximum routed interfaces: 8,000

- Maximum unique combinations: 32,000 (source, group, receiver)

Flows

Flows provides deep insights at a flow level giving details such as average latency and packet drop indicator. It also raises anomalies when the latency of the flows increase or when packets get dropped because of congestion or forwarding errors.

Each flow has a packet counter representing the number of packets entering the ASIC for that flow over a period of time. This period of time is called aggregation interval. There are several points where flow statistics for a given flow can be aggregated. Aggregation can happen in the ASIC, switch software, and server software.

The Flows section of Nexus Dashboard displays the telemetry information collected from various devices in the fabric that were added to the fabric.

For details on Flow Telemetry support for Cisco Nexus series switches and line cards, see the "Compatibility Information" section in the *Nexus Dashboard Release Notes*.

Flows hardware requirements

For details on Flow Telemetry support for Cisco Nexus platform switches, see the "Compatibility Information" section in the *Nexus Dashboard Release Notes*.

Flows guidelines and limitations for NX-OS fabrics

For details on Flow Telemetry hardware support, see the "Compatibility Information" section in the *Nexus Dashboard Release Notes*.

- Output port information for outgoing traffic from N9K-C93180YC-EX, N9K-C93108TC-EX, and N9K-C93180LC-EX line cards will not be displayed.
- Flows does not support multicast traffic. The access list must be provisioned to exclude the multicast traffic flows.
- A maximum of 63 VRFs are supported on flow telemetry nodes.
- The number of anomalies in the **Fabric Overview** dashboard will not match the number of anomalies in the flow browse page. The Fabric Dashboard contains the total anomaly count for the time range you selected. The flow records are not aggregated in the flow browse view, where multiple flow records can point to the same anomaly entry.
- The L3-VNI flows show as L2-VNI flows when the VXLAN flow is dropped in the ingress node. When VXLAN packets are dropped in the first-hop, the exported VXLAN flow telemetry records will indicate the drop. However, they don't carry the VNI information in it. The Ingress interface from the flow telemetry export along with the VRF associated with the interface, does not deduce if the flow is L2-VNI or L3-VNI. In this case Nexus Dashboard associates the L2-VNI for the flow.
- When a VXLAN encapsulated packet enters an Cisco N9500-EX switch and feature overlay (EVPN) is configured, the packet will be treated like a VXLAN transit node packet. Also the ingress interface and egress interface are set as zeros in the flow telemetry export. The ingress and egress interfaces are needed to consider this record for flows. The limitation on these switches results in the Cisco N9500-EX switch not being considered in the path stitching and correlation if the switch is in ingress, transit, or egress direction. Cisco N9500-EX switches will be treated like a transit node for an overlay packet.

- For Nexus Dashboard to work in VXLAN deployments, you must have symmetric configuration on the switches involved in the overlay. This enables Nexus Dashboard to correlate and stitch the overlay flows. When such a symmetric configuration is not present, the VXLAN feature and forwarding will work, but Nexus Dashboard will not stitch the flows correctly. See the following examples to understand what is meant by symmetric configuration on switches:
 - For Layer 2 VXLAN VNI cases: If vlan-x is mapped VNI-A in PE1, then the same vlan-x must be mapped to VNI-A in PE2, where PE1 and PE2 are VTEP endpoints for the Layer 2 overlay.
 - For Layer 3 VXLAN VNI cases: If SVI-x is mapped to VRF-A mapped VNI-P on PE1, then the same SVI-x must be mapped to VRF-A mapped VNI-P in PE2, where PE1 and PE2 are VTEP endpoints for the Layer 3 overlay.
- The ingress and VRF information will not be shown for all interfaces which use the flow telemetry 'tenant-id' for encoding the logical interface ID, as this ID will be used for 'overlay-id'. It's not possible to derive the logical interface (SVI with trunk port, sub interface, SVI with trunk and port channel) and get the VRF associated with it. This results in the flow browse page and details page not showing the ingress and egress VRFs.
- On the Cisco N9500-EX switches connected to VPC pair, the current design limits in identifying the ingress leaf nodes between VPC pairs causing the loss of flow in Nexus Dashboard.
- Large volumes of anomalies in the indices degrade flow database write performance and cause Kafka lag on the 350 nodes that support software and flow telemetry. This lag prevents Nexus Dashboard from displaying complete data.
- Flows information is retained for 7 days or until flow database reaches 80%, which ever happens first, then older flows information is deleted from the database.
- Flow telemetry and flow telemetry events will not export **drop bit** if there is an egress ACL drop in Cisco Nexus FX switches.
- For Nexus Dashboard to receive Flow Telemetry data, the TCAM region for **ing-netflow** must be set to 512. See [N9000 TCAM Carving](#).
- For flows, if the time range you have selected is greater than 6 hours, the data may not get displayed. Select a time range that is less than or equal to 6 hours.
- When the scale limit for anomalies is reached, some of the unhealthy flows may not be displayed as anomalies in the Flow Record Details page. A system issue is raised when this condition happens. Navigate to **Admin > System Settings > System Issues** to view the system issue.
- For Layer 4 to Layer 7 Services only intra VRF is supported for flows. For Layer 4 to Layer 7 Services inter VRF is not supported for flows.
- Multicast traffic is not supported by classic flow telemetry. Classic flow telemetry collects per-flow statistics (such as anomaly score, latency, and packet drop indicators) for unicast traffic only. To monitor multicast traffic, use the multicast telemetry flow path visibility feature instead. For more information, see [Multicast traffic monitoring overview](#).
- When traffic flows through sub-interfaces, in the **Flow Record Details** page, sub-interface is displayed only in the ingress direction in the Flow Path area. In the egress direction the parent interface is displayed.
- For unicast routes in Flow Telemetry, Layer 4 to Layer 7 Services traffic path visibility is supported on Cisco N9300 -GX2 and -FX3 platform switches.

- Configuring Netflow on a switch before onboarding a Standalone NX-OS fabric on Nexus Dashboard and then enabling Flow Telemetry on the corresponding fabric from Nexus Dashboard could lead to anomalies being generated for incorrect configurations being pushed to the switch.

To resolve the issue, perform the following steps:

1. Remove switch from the Standalone NX-OS fabric.
 2. Use the command **no feature netflow** to remove Netflow from the switch.
 3. Add the switch back to the Standalone NX-OS fabric.
 4. Enable Flow Telemetry.
- When Flow Telemetry is disabled while one of the switches is unreachable, the fabric goes into the **Disable Failed** state. This is expected behavior. Following this condition, when the switch becomes available and you enable Flow Telemetry, the ACL configurations get corrupted.

We recommend either of these workarounds:

- When this condition occurs and the fabric is in a **Disable Failed** state, retry the switch that is in this state. This will trigger the ACL to unconfigure successfully. Thereafter, when you enable the fabric, start a **disable acl job** to clean any existing ACLs in the switch.
- If you have already enabled Flow Telemetry without cleaning the existing ACLs in the switch, then perform these actions:
 1. Disable Flow Telemetry and then enable it.
 2. Remove the switch that has the issue from the DCNM fabric and add it back in the fabric.
- If an NX-OS fabric contains a mix of Cloud Scale (CS) and Silicon One (E100 or G200) switches, PTP is not supported and flow telemetry cannot be enabled.

View flows

The **Flows** page displays telemetry information collected from various devices in an online fabric. The Flows records let the user visualize the flows per fabric. For a particular fabric, you can view flows by the following metrics.

- Anomaly Score
- Drop Score
- Latency Score
- Packet Drop Indicator
- Average Latency

The flows engine also runs machine-learning algorithms on the behavior of the flows to raise anomalies in the behavior, such as average latency and packet drop indicator. The graph represents the anomalies in the behavior over a period of time.

Flow telemetry and analytics gives in-depth visibility of the data plane. Flows collects the flow records streamed from the nodes and converts to understandable EPG-based flow records. Top nodes by flow anomalies displays the nodes in the network with the most anomalies.



To view flow details, you must first enable flows.

1. Click **Manage > Fabrics**, then click on the appropriate fabric.
2. Click **Actions > Edit Fabric Settings**.
3. Under **General**, verify that the **Telemetry** option is enabled in the **Enabled features** area.
4. Under **Telemetry**, verify that the **Flow telemetry** option is enabled in the **Mode** area.

The **Flows** page provides deep visibility into your network traffic, allowing you to monitor both unicast and multicast flows. For more information see, [View unicast flows](#) and [View multicast flows](#).

View unicast flows

You can view unicast flows to monitor traffic patterns, identify anomalies, and troubleshoot performance issues across your fabric.

Follow these steps to view unicast flows.

1. Navigate to the **Fabrics** page.

Go to **Manage > Fabrics**.

2. Click the fabric you configured to collect multicast data.

The **Fabric Overview** page displays.

3. Click the **Connectivity** tab and then click the **Flows** tab.
4. Click the **Unicast** vertical tab.
5. Choose a time range from the time-range selector.
6. In **Flows by Anomaly score** area, choose an option from the **Site Flows by** drop-down list.

Choose from the following options.

- **Anomaly score**—The score reflects the number of detected anomalies logged in the database. A higher score indicates a greater number of anomalies associated with the flow, helping you prioritize investigation of problematic flows.
- **Drop score**—The score quantifies the severity of packet drops observed for a flow. Nexus Dashboard calculates the score based on the frequency and volume of dropped packets detected through flow records received from the switch. A higher drop score indicates more significant packet loss, which may impact application performance.
- **Latency score**—The score represents the severity of latency experienced by a flow as it traverses the fabric. Nexus Dashboard derives the score from the measured end-to-end latency values and flags flows that exceed expected latency thresholds. A higher latency score indicates greater delay and potential performance degradation.
- **Packet drop indicator**—Nexus Dashboard analyzes flow records to detect packet drops. The primary detection method uses the drop bit received from the switch within the flow records, enabling rapid identification of flows experiencing loss along the forwarding path.
- **Average latency**—The metric measures the time a packet takes to traverse from source to destination within the fabric. To ensure accurate fabric latency measurement, synchronize all

nodes to a uniform time source (for example, through PTP or NTP).

7. Review the **Flows** table to see detailed information for each flow, including anomaly score, latency score, drop score, flow record time, nodes, flow type, protocol, latency, and packet drop indicator.



To narrow the results, enter a criteria in the **Filter by attributes field**. To sort the results, click any column header.

Layer 4 to Layer 7 services traffic path visibility

You have expanded visibility in the Flow Path to Layer 4 to Layer 7 services external devices such as firewalls. Nexus Dashboard tracks the end-to-end flow across the service chain in real-time and helps locate data plane issues across the device silos. A non-NAT environment across all third-party vendors is supported.

For Layer 4 to Layer 7 services traffic path visibility, your flow telemetry must be enabled and the appropriate rules must be configured.

1. Click **Manage > Fabrics**, then click on your fabric.
2. Click **Actions > Edit Fabric Settings**.
3. Under **General**, verify that the **Telemetry** option is enabled in the **Enabled features** area.
4. Under **Telemetry > Flow Collection**, verify that the **Flow telemetry** option is enabled in the **Mode** area.

Based on your rules, if the flow is passing through Policy Based Redirects (for example, a firewall), it will display that information in the flow path.

View Layer 4 to Layer 7 services traffic path visibility

Follow these steps to view Layer 4 to Layer 7 services traffic path visibility.

1. Navigate to **Manage > Fabrics**.
2. Select **Online Fabrics** from the drop-down list.
3. Click a fabric name to view the fabric details.
4. Navigate to **Connectivity > Flows**.
5. Select a time range.
6. In the flows table, click Record Time to view the flow record details.

In the **Path** area, a graphical flow path will display the end-to-end information, from source to destination, and it will also identify the firewall in the path if a firewall is present. The graph also captures the end-to-end flow path network latency that is occurring. In the graph, if there are any anomalies, a red dot is displayed next to the symbol for the leaf switch or the spine switch.

7. Click **Anomalies** in the **Flow Details** page to view further details related to the anomaly.



In the current release firewalls are not supported for anomalies.

Guidelines and limitations for Layer 4 to Layer 7 services traffic path visibility

These guidelines and limitations are for Layer 4 to Layer 7 services traffic path visibility.

- This feature is currently recommended only if policy-based redirect can be configured using the service graph for a Cisco ACI fabric or Layer 4 to Layer 7 services for an NX-OS fabric.
- In the current release, firewalls are not supported for anomalies.
- In the current release, the latency information that is being displayed is the network latency, and it does not capture the latency that is occurring in the firewall.
- In the current release, NAT is not supported.
- This feature is currently supported if you use these switches:
 - Cisco N9300-FX Platform Switches
 - Cisco N9300-FX2 Platform Switches
 - Cisco N9300-FX3 Platform Switches
 - Cisco N9300-GX Platform Switches
- The flow path displays an **unknown** node for L3Out-based firewalls or Layer 4 to Layer 7 services. In these configurations, the APIC creates an implicit VRF for traffic forwarding, which prevents switches from generating flow telemetry records.
- In a Cisco ACI fabric, when the Layer 4 to Layer 7 services traffic is forwarding in Layer 2 (bridged mode), the flow analytics support is limited. Ingress and egress nodes detection is not accurate, and path summary is not available.
- The service node displays as **unknown** for Layer 4 to Layer 7 traffic path visibility.
- Without a service graph for Layer 4 to Layer 7 services, if the client > service node is VRF_A and the service node > server is VRF_B, the paths will be recorded as separate flows as there is no common or single contract to stitch the flows.
- Load balancers are not supported.

Flow telemetry events

Flow telemetry events are enabled implicitly when flow telemetry is enabled, and on NX-OS fabrics flow rules are configured. The flow telemetry enables triggering events when a configured rule is met, where packets are exported to the collector for analysis.

Flow telemetry events enhance and complement current flows in Nexus Dashboard. They enrich anomaly generation for flow telemetry and flow telemetry events.

It monitors security, performance, and troubleshooting. This is achieved using the periodic flow table event records exported every second.

The data export to Nexus Dashboard is done directly from the hardware without control plane needing to handle the data. Statistics are assembled as a packet with a configurable MTU size and a defined header. These packets are sent as in-band traffic from the fabric. Headers are configured by software, and packets streamed are UDP packets.

When flow telemetry is available for a triggered flow telemetry event, then you can navigate to flow details page for aggregated information. These events are based on the following drop events:

- **Cisco ACL Drop**—In an NX-OS fabric, when packet hits **sup-tcam** rules and the rule is to drop the packet, the dropped packet is counted as ACL_Drop and it will increment the forward drop counter. When this occurred, it usually means the packet is about to be forwarded against basic Cisco ACI forwarding principals. The **sup-tcam** rules are mainly to handle some exceptions or some of control plane traffic and not intended to be checked or monitored by users.
- **Buffer Drop**—When the switch receives a frame and there are no buffer credits available for either ingress or egress interface, the frame is dropped with buffer. This typically hints at a congestion in the network. The link that is showing the fault could be full or the link containing the destination may be congested. In this case a buffer drop is reported in flow telemetry events.
- **Forward Drop**—The packets that are dropped on the LookUp block (LU) of the Cisco ASIC. In a LU block a packet forwarding decision is made based on the packet header information. If the packet is dropped, forward drop is counted. There may be a variety of reasons when forward drop is counted.
- In NX-OS fabrics, a **Forward Drop** (for example, due to a route miss or FIB exception) may trigger an **acl_drop** event at the ASIC level, especially when the system is not in **cfg_aci_mode**. This occurs because FIB exceptions redirect packets to supervisor ACLs with control plane policers, and exceeding these policers causes the **acl_drop** event, which flow telemetry may internally classify as a policy drop."
- Nexus Dashboard does not typically raise a dedicated policy drop anomaly for certain instances; instead, the 'Packet drop indicator' for the affected flow increments. Therefore, if you observe an increased 'Packet drop indicator' for traffic that appears to have routing or forwarding problems, the underlying issue might be a forwarding drop, despite the internal ASIC classification.

When a TCP retransmission happens for a flow, it indicates potential packet loss or forwarding issues within the fabric.



These specific behaviors and anomaly definitions are applicable only to Cisco ACI fabrics and not apply to NX-OS fabrics or other fabric types.

Flow telemetry events compared to flow telemetry

- The flow telemetry event packets are exported only when configured events occur, whereas flow telemetry packets are streamed continuously.
- The flow telemetry events are captured for all traffic, where as flow telemetry is captured for filtered traffic.
- The total number of collectors between flow telemetry and flow telemetry events is 256.

Guidelines and limitations for flow telemetry events

- In a standalone NX-OS fabric, flow telemetry event anomalies are aggregated. For example, a packet drop anomaly occurred from time T0 to T1. No packet drop anomaly occurred from time T1 to T2. Another packet drop anomaly occurred from time T2 to T3. Although there is no anomaly from T1 to T2, the time stamp for the aggregated packet drop anomalies is from T0 to T3.
- The flow telemetry events do not report policing drop anomalies in Nexus Dashboard, when the egress data plane policer is configured on front-panel ports and there is traffic drop.
- To export flow telemetry events on FX platform switches, you must configure flow telemetry filters. In a Cisco ACI fabric, starting with Cisco ACI-Mode Switch release 16.0(3), FX switches

export flow telemetry events to indicate only buffer drops experienced by flows without the need to configure flow telemetry filters.

- Standalone NX-OS fabrics do not support TCP packet RTO anomalies.

View multicast flows

You can view multicast flows to monitor active multicast traffic in the fabric, analyze source and receiver behavior, and visualize the end-to-end flow path from the FHR to the LHR.

Follow these steps to view multicast flows.

1. Navigate to the **Fabrics** page.

Go to **Manage > Fabrics**.

2. Click the fabric you configured to collect multicast data.

The **Fabric Overview** page displays.

3. Click the **Connectivity** tab and then click the **Flows** tab.

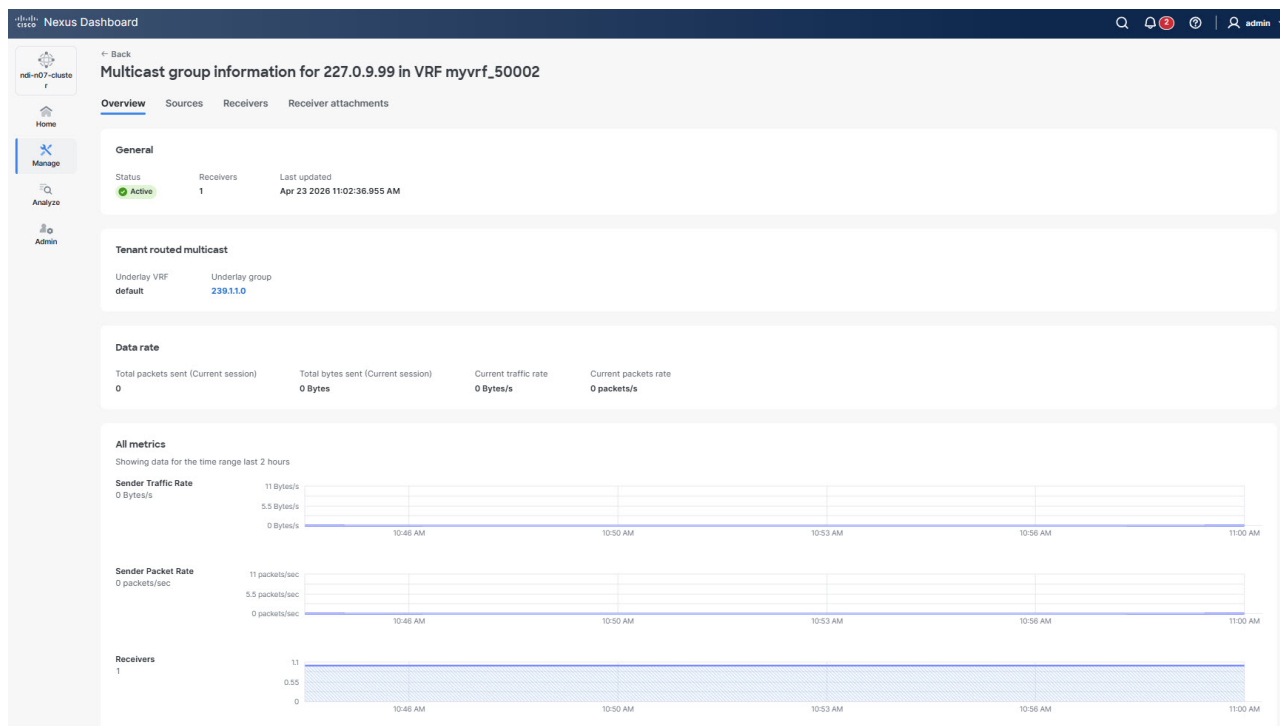
4. Click the **Multicast** vertical tab.

5. Review the flow table for the following details.

- **Status:** Indicates if the flow is **Active/Inactive**.
- **Multicast group:** The multicast IP address.
- **Tenant and VRF:** The tenant and VRF context for the flow.
- **Source, Connected switch, and Connected interface:** Identifies the source (S) and its attachment point in the fabric.
- **Source traffic rate:** The current bandwidth utilization for the flow.

6. Click a multicast group entry in the **Flows** table to view detailed source and receiver information.

The multicast group information page displays.



7. Click the following tabs to monitor specific multicast metrics.

- **Overview**—displays the current status (**Active/Inactive**), tenant routed multicast details including underlay VRF and underlay group, data rate information (packets sent and bytes sent), and the **All metrics** section with trend graphs for **Sender traffic rate**, **Sender packet rate**, and **Receivers**.
- **Sources**—displays each source contributing to the group, including source IP address, connected switch name, and ingress interface.
- **Receivers**— displays each receiver for the group, including receiver IP address, connected switch, interface, and join status.

8. From the multicast group information page, click the **Sources** tab.

9. From the **Sources** table, click the **Source IP address**.

The source information page displays.

Nexus Dashboard

bl11-nxos-cluster

Home

Manage

Analyze

Admin

← Back

Source 192.168.101.2 in multicast group 225.0.0.1 in tenant Tenant-A, VRF vrf101

View in topology

Overview Receivers

General

Status	Receivers	Last updated	VRF
Active	1 →	May 27 2026 12:19:33.840 PM	vrf101

Tenant routed multicast

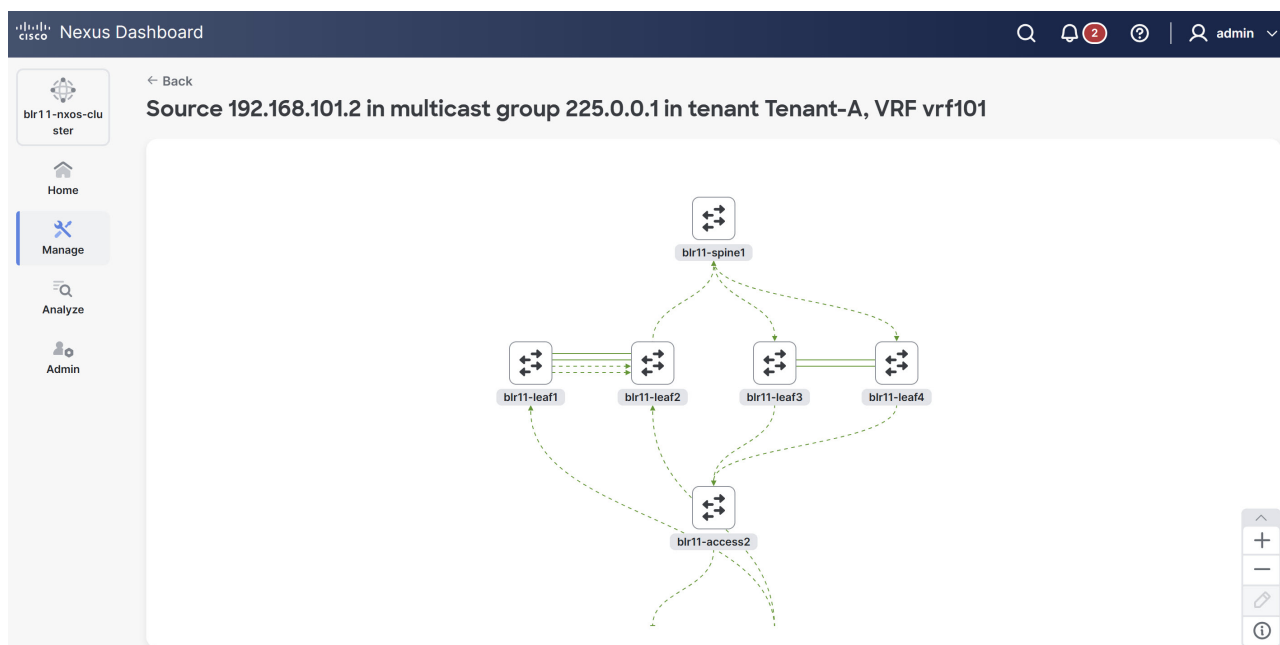
Underlay VRF	Underlay group
default	239.1.1.0

Data rate

Total packets sent	Total bytes sent	Source traffic rate (Last 5 min)	Source packet rate (Last 5 min)
136.28 M	16.25 GB	1.03 mbps →	1.00 K packets/s →

10. Click **View in topology** to visualize the end-to-end multicast flow path.

The topology view displays the complete forwarding path from the FHR through intermediate switches to each LHR, along with per-hop flow statistics.



gasurian: commenting out the following sections as they are applicable only for IPFM fabrics. We gave an article for Working with connectivity in IPFM fabrics and it will be included to the collection page/Help center from ND 4.3

Virtual Infrastructure

To work with virtual infrastructure connectivity in a fabric:

1. [Navigate to the Connectivity page.](#)
2. Click the **Virtual Infrastructure** tab.

A page showing the information for previously-configured virtual infrastructure connectivity appears.

View virtual machine VMs

The **Actions** drop-down list displays only for fabrics with security groups enabled. For fabrics with security groups enabled, you need to perform a **Recalculate and deploy** operation only once. For every **Actions > Set group ID** or **Actions > Reset Group ID** operation, **Recalculate and deploy** is not needed. For more information on security groups, see [Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric](#).

Follow these steps to navigate to the **Virtual Machine VMs** page.

1. [Navigate to the Connectivity page](#).
2. Click the **Virtual Infrastructure** tab.
3. Click the **Virtual Machine VMs** subtab.

The **Virtual Machine VMs** subtab displays detailed information for VMs.

Field	Description
VM Name	Specifies the name of the virtual machine.
IP Address	Specifies the IP address of the virtual machine.
MAC Address	Specifies the MAC address of the virtual machine.
VLAN	Specifies the VLAN associated with the virtual machine.
Network	Specifies the network associated with the virtual machine.
VRF	Specifies the VRF associated with the virtual machine.
Switch	Specifies the switch connected to the virtual machine.
Switch interface	Specifies the switch interface connected to virtual machine.
State	Specifies the state of the VM adapter.

You can search and filter VMs by using the **Filter by attributes** search field.

These are the options available through the **Actions** menu in the **Virtual Machine VMs** page:

Field	Description
Set group ID	Used to associate a VM with a security group. 1. Choose the VM that you want to associate with a security group. 2. Click Actions > Set group ID. 3. Choose an existing security group, or click Add Security Group. See the section "Create a security group" in Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric for more information. 4. Click Submit. This action configures an IP selector for the VRF on the switches.
Remove group ID	Used to remove an association with a security group on a VM. To remove an association with a security group on a VM, choose that VM and click Actions > Remove group ID.

You can search and filter VMs by using the **Filter by attributes** search field.

View Kubernetes pods

The **Actions** drop-down list displays only for fabrics with security groups enabled. For fabrics with security groups enabled, you need to perform a **Recalculate and deploy** operation only once. For every **Actions > Set group ID** or **Actions > Reset Group ID** operation, **Recalculate and deploy** is not needed. For more information on security groups, see [Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric](#).

Follow these steps to view Kubernetes pods on the **Fabrics** page.

1. [Navigate to the Connectivity page](#).
2. Click the **Virtual Infrastructure** tab.
3. Click the **Kubernetes Pods** subtab.

You can search and filter Kubernetes pods by using the **filter by attributes** search field.

The **Kubernetes Pods** subtab displays detailed information for Kubernetes pods.

Field	Description
Pod Name	Specifies the name of the Kubernetes pod.
Pod IP	Displays the IP address of the Kubernetes pod.
Phase	Specifies the phase (state) of the pod.
Reason	Specifies the reason.
Applications	Specifies the applications of the pod.
Namespace	Specifies the namespace of the pod.

Field	Description
Node Name	Specifies the node name of the pod.
Node IP	Specifies the node IP address.
Cluster Type	Displays the type of cluster.
Physical NIC	Displays the physical NIC of the node.
Physical Switch	Specifies the physical switch connected to cluster node.
Switch Interface	Specifies the switch interface connected to cluster node.
Cluster Name	Specifies the name of the cluster.
Port Channel	Specifies the port channel (if cluster node is connected to a VPC).
VLAN	Specifies the VLAN.
Fabric	Specifies the fabric name.

Copyright

THE SPECIFICATIONS AND INFORMATION REGARDING THE PRODUCTS IN THIS MANUAL ARE SUBJECT TO CHANGE WITHOUT NOTICE. ALL STATEMENTS, INFORMATION, AND RECOMMENDATIONS IN THIS MANUAL ARE BELIEVED TO BE ACCURATE BUT ARE PRESENTED WITHOUT WARRANTY OF ANY KIND, EXPRESS OR IMPLIED. USERS MUST TAKE FULL RESPONSIBILITY FOR THEIR APPLICATION OF ANY PRODUCTS.

THE SOFTWARE LICENSE AND LIMITED WARRANTY FOR THE ACCOMPANYING PRODUCT ARE SET FORTH IN THE INFORMATION PACKET THAT SHIPPED WITH THE PRODUCT AND ARE INCORPORATED HEREIN BY THIS REFERENCE. IF YOU ARE UNABLE TO LOCATE THE SOFTWARE LICENSE OR LIMITED WARRANTY, CONTACT YOUR CISCO REPRESENTATIVE FOR A COPY.

The Cisco implementation of TCP header compression is an adaptation of a program developed by the University of California, Berkeley (UCB) as part of UCB's public domain version of the UNIX operating system. All rights reserved. Copyright © 1981, Regents of the University of California.

NOTWITHSTANDING ANY OTHER WARRANTY HEREIN, ALL DOCUMENT FILES AND SOFTWARE OF THESE SUPPLIERS ARE PROVIDED "AS IS" WITH ALL FAULTS. CISCO AND THE ABOVE-NAMED SUPPLIERS DISCLAIM ALL WARRANTIES, EXPRESSED OR IMPLIED, INCLUDING, WITHOUT LIMITATION, THOSE OF MERCHANTABILITY, FITNESS FOR A PARTICULAR PURPOSE AND NONINFRINGEMENT OR ARISING FROM A COURSE OF DEALING, USAGE, OR TRADE PRACTICE.

IN NO EVENT SHALL CISCO OR ITS SUPPLIERS BE LIABLE FOR ANY INDIRECT, SPECIAL, CONSEQUENTIAL, OR INCIDENTAL DAMAGES, INCLUDING, WITHOUT LIMITATION, LOST PROFITS OR LOSS OR DAMAGE TO DATA ARISING OUT OF THE USE OR INABILITY TO USE THIS MANUAL, EVEN IF CISCO OR ITS SUPPLIERS HAVE BEEN ADVISED OF THE POSSIBILITY OF SUCH DAMAGES.

Any Internet Protocol (IP) addresses and phone numbers used in this document are not intended to be actual addresses and phone numbers. Any examples, command display output, network topology diagrams, and other figures included in the document are shown for illustrative purposes only. Any use of actual IP addresses or phone numbers in illustrative content is unintentional and coincidental.

The documentation set for this product strives to use bias-free language. For the purposes of this documentation set, bias-free is defined as language that does not imply discrimination based on age, disability, gender, racial identity, ethnic identity, sexual orientation, socioeconomic status, and intersectionality. Exceptions may be present in the documentation due to language that is hardcoded in the user interfaces of the product software, language used based on RFP documentation, or language that is used by a referenced third-party product.

Cisco and the Cisco logo are trademarks or registered trademarks of Cisco and/or its affiliates in the U.S. and other countries. To view a list of Cisco trademarks, go to this URL: <https://www.cisco.com/go/trademarks>. Third-party trademarks mentioned are the property of their respective owners. The use of the word partner does not imply a partnership relationship between Cisco and any other company. (1110R)

© 2017–2026 Cisco Systems, Inc. All rights reserved.

Americas Headquarters

Cisco Systems, Inc.

170 West Tasman Drive

San Jose, CA 95134-1706

USA

<https://www.cisco.com>

Tel: 408 526-4000

800 553-NETS (6387)

Fax: 408 527-0883