



Editing AI Data Center VXLAN Fabric Settings, Release 4.3.1

Table of Contents

New and changed information	1
Editing AI Data Center VXLAN fabric settings	2
General	3
Fabric Management	5
General Parameters	5
Replication	8
vPC	10
Protocols	12
Security	16
Advanced	19
Freeform	26
Resources	26
Manageability	32
Bootstrap	33
Configuration Backup	35
Flow Monitor	36
Telemetry	39
Edit configuration settings	40
Edit NAS settings	40
Disable Telemetry	41
Perform force disable telemetry on your fabric	41
NAS	41
Guidelines and limitations for network attached storage	42
Add network attached storage to export flow records	42
Add NAS to Nexus Dashboard	42
Add the onboarded NAS to Nexus Dashboard	43
Flow collection	45
Understanding flow telemetry	45
Flow telemetry guidelines and limitations	45
Configure flows	47
Monitor the subnet for flow telemetry	50
Understanding Netflow	50
Understanding Netflow types	51
Netflow guidelines and limitations	51
Configure Netflow	52
Understanding sFlow	53
Guidelines and limitations for sFlow	53
Configure sFlow telemetry	53
External streaming	55
Configure external streaming settings	55
Guidelines and limitations	56

Guidelines and limitations in NX-OS fabrics	56
Email	56
Message bus	58
Add Kafka broker configuration	58
Configure Kafka exports in fabric settings	60
Anomalies	62
Advisories	63
Statistics	63
Faults	64
Audit Logs	64
Syslog	64
Guidelines and limitations for syslog	65
Add syslog server configuration	65
Configure syslog to enable exporting anomalies data to a syslog server	65
Additional settings	67
VXLAN EVPN fabrics provisioning	67
Guidelines for VXLAN BGP EVPN fabrics provisioning	68
AI fabric management	70
AI settings	71
Layer 3 VNI without VLAN	73
Guidelines and limitations for Layer 3 VNI without VLAN	74
Precision Time Protocol for Data Center VXLAN EVPN fabrics	74
AI QoS classification and queuing policies	75
Understanding AI QoS classification and queuing policies	76
Guidelines and limitations for AI QoS classification and queuing policies	76
Configure AI QoS classification and queuing policies	77
Create a policy using the custom QoS templates	78
MACsec support in Data Center VXLAN EVPN and BGP fabrics	79
Guidelines	79
Enable MACsec	79
Disable MACsec	84
Provisioning VXLAN EVPN Fabric with IGP underlay	85
Create a VXLAN EVPN fabric with IPv4 underlay	85
Create a VXLAN EVPN fabric with IPv6 underlay	85
Add switches	89
Assigning Switch Roles	89
vPC fabric peering	90
Overlay mode	95
Managing a brownfield VXLAN BGP EVPN fabric	99
Prerequisites	100
Guidelines and limitations	100
Fabric topology overview	101
Nexus Dashboard brownfield deployment tasks	102

Configuration profiles support for brownfield migration	110
Manually add PIM-BIDIR configuration for leaf or spine post brownfield migration	111
Migrate interconnected VXLAN fabrics with border gateway switches	111
Configuring a VXLANv6 fabric	113
Guidelines and limitations for an IPv6 underlay	113
Create a VXLAN EVPN fabric with IPv6 underlay	114
Configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6	118
Overview of VXLAN EVPN to SR-MPLS and MPLS LDP interconnection	121
Supported platforms and configurations	121
Inter-fabric connections for MPLS handoff	122
VXLAN MPLS topology	122
Configuration tasks for VXLAN MPLS handoff	124
Editing fabric settings for MPLS handoff	125
Create an underlay inter-fabric connection	126
Create an overlay inter-fabric connection	127
Deploy VRFs	128
Change the routing protocol and MPLS settings	129
Overview of Tenant Routed Multicast	130
Guidelines and limitations	130
Overview of Tenant Routed Multicast with interconnected VXLAN fabrics	131
Operations for Tenant Routed Multicast with interconnected VXLAN fabrics	131
Configure Tenant Routed Multicast for a single site	131
Configure Tenant Routed Multicast with interconnected VXLAN fabrics	134
Understanding enhanced analytics for AI fabrics	135
AI endpoint and topology discovery	136
Prerequisites	137
Guidelines and limitations: Analytics for AI fabrics	137
Add Slurm as an integration	139
DCGM Exporter	139
Configure DOCA Telemetry Service and Fluent Bit	140
GPU/NIC Topology Fleet Collector	152
Collect data in secured mode	156
Associate a fabric with an AI cluster	158
Associate a switch with a scalable unit	159
View information on AI resources	160
View AI job information	168
Copyright	187

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	Enhanced border node functions and routing isolation for multi-tier VXLAN fabrics	Beginning with Nexus Dashboard 4.3.1, you can use unique ASNs on applicable border switches in AI VXLAN eBGP fabrics, use per-VRF per-VTEP unique loopback IP auto-provisioning in AI VXLAN eBGP fabrics, allow border and border gateway connections to spine switches in the presence of super spines, and automatically generate EVPN overlay peering policies for AI VXLAN eBGP fabrics. For more information, refer to General Parameters , Replication , and Resources .
Nexus Dashboard 4.3.1	Support for Cisco N9348Y12C-SE1 switches	Starting with Nexus Dashboard 4.3.1, Nexus Dashboard adds support for Cisco N9348Y12C-SE1 switches running NX-OS 10.6(3) release on all LAN fabric types. For more information, refer to Add switches .
Nexus Dashboard 4.3.1	Cisco Nexus 9000 Series name change	Cisco Nexus 9000 Series is now the Cisco N9000 Series.

Editing AI Data Center VXLAN fabric settings

An **AI Data Center VXLAN** fabric is a type of fabric that is used for a VXLAN EVPN deployment with Cisco N9000 and Cisco Nexus 3000 switches.

When you first create an AI Data Center VXLAN fabric using the procedures provided in [Creating Fabrics and Fabric Groups](#), the standard workflow allows you to create a fabric using the bare minimum settings so that you are able to create a fabric quickly and easily. Use the procedures in this article to make more detailed configurations for your AI Data Center VXLAN fabric.



You can create an AI Data Center VXLAN EVPN fabric with an IPv6 underlay. The IPv6 underlay is supported for VXLAN EVPN templates. For information about the IPv6 underlay, see [Configuring a VXLANv6 fabric](#) and [Configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6](#).

1. Navigate to the main Fabrics window:

Manage > Fabrics

2. Locate the AI Data Center VXLAN fabric that you want to edit.

AI Data center VXLAN fabrics are shown with **AI Data center VXLAN EVPN** in the **Type** column.

3. Click the circle next to the AI Data Center VXLAN fabric that you want to edit to select that fabric, then click **Actions > Edit Fabric Settings**.

The **Edit *fabric_name* Settings** page appears.

4. Click the appropriate tab to edit these settings for the fabric:



If you're creating a standalone fabric as a potential member fabric of a VXLAN fabric group (used for provisioning overlay networks for fabrics that are connected), see [Creating Fabrics and Fabric Groups](#) before creating the member fabric.

- [General](#)
- [Fabric Management](#)
- [Telemetry](#) (if the **Telemetry** feature is enabled for the fabric)
- [AI settings](#)

General

Use the information in this section to edit the settings in the **General** page for your AI Data Center VXLAN fabric.

Change the general parameters that you configured previously for the AI Data Center VXLAN fabric, if necessary, or click another tab to leave these settings unchanged.

Fabric type	Description
Name	The name for the fabric. This field is not editable.
Type	The fabric type for this fabric. This field is not editable.
Location	Choose the location for the fabric.
Routing Protocol	Displays the type of routing protocol chosen: ▪ iBGP: Interior Border Gateway Protocol ▪ eBGP: Exterior Border Gateway Protocol
BGP ASN for Spines	Enter the BGP autonomous system number (ASN) for the fabric's spine switches. Options are: ▪ 400G: Enable QoS queuing policies for an interface speed of 400 Gb. ▪ 100G: Enable QoS queuing policies for an interface speed of 100 Gb. ▪ 25G: Enable QoS queuing policies for an interface speed of 25 Gb.
License tier	Choose the licensing tier for the fabric: ▪ Essentials ▪ Advantage ▪ Premier Click on the information icon (i) next to License tier to see what functionality is enabled for each license tier.
Enabled features	Check the box to enable Telemetry for the fabric. This is the equivalent of enabling the Nexus Dashboard Insights service in previous releases.
Telemetry collection	This option becomes available if you choose to enable Telemetry in the Enable features field above. Choose either Out-of-band or In-band for telemetry collection.
Telemetry streaming	This option becomes available if you choose to enable Telemetry in the Enable features field above. Choose either IPv4 or IPv6 for telemetry streaming.
Security domain	Choose the security domain for the fabric.



Deploying a new AI VXLAN fabric in Nexus Dashboard release 4.3.1 or upgrading an

existing AI VXLAN fabric from release 4.2.1 to 4.3.1 on the Premier licensing tier updates your leaf switch topology view. A single Scalable Unit icon replaces individual leaf switch icons and initially displays an "Unassigned" status. Nexus Dashboard performs this update to integrate the AI resource analytics enhancements introduced in this release.

- To view individual leaf switches: Double-click the Scalable Unit icon to expand the view and display all switches contained within that unit.
- To rename the scalable unit: Associate the switches with a scalable unit. Nexus Dashboard labels new units as "Unassigned" by default. For more information, see [Associate a switch with a scalable unit](#).

Fabric Management

Use the information in this section to edit the settings in the **Fabric Management** window for your AI Data Center VXLAN fabric. The tabs and their fields in the screen are explained in the following sections. The fabric-level parameters are included in these tabs.

- [General Parameters](#)
- [Replication](#)
- [vPC](#)
- [Protocols](#)
- [Security](#)
- [Advanced](#)
- [Freeform](#)
- [Resources](#)
- [Manageability](#)
- [Bootstrap](#)
- [Configuration Backup](#)
- [Flow Monitor](#)

General Parameters

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

Field	Description
Enable IPv6 Underlay	Enable the IPv6 underlay feature. For more information, see the section "Configuring a VXLANv6 Fabric" in Configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6 .
Enable IPv6 Link-Local Address	Enables the IPv6 link-local address.
Fabric Interface Numbering	Specifies whether you want to use point-to-point (p2p) or unnumbered networks.
Underlay Subnet IP Mask	Specifies the subnet mask for the fabric interface IP addresses.
Underlay Subnet IPv6 Mask	Specifies the subnet mask for the fabric interface IPv6 addresses.
Underlay Routing Protocol	The IGP used in the fabric, OSPF or IS-IS.
ASN Auto Allocation	Enable automated BGP ASN allocation for leaf switches, border devices, and border gateways in Multi-AS mode. Nexus Dashboard assigns unique ASNs from the defined pool. This field applies only when the underlay and overlay routing protocol is eBGP.

BGP ASN Range	Specifies the BGP ASN range for automatic allocation. This field is available when ASN Auto Allocation is enabled. If you do not provide a value, the system automatically generates a range. Use a unique BGP ASN range for each eBGP fabric.
Allow Unique ASN on Border	Select the check box to allow supported border switches and border gateways to use different ASNs when the fabric is configured to use the Same-Tier-AS ASN assignment model for border-class switches. When this option is enabled, the fabric-wide border ASN setting is unavailable. You must configure the ASN for each applicable switch individually by using the leaf_bgp_asn policy.  When the Allow Unique ASN on Border option is enabled, you must manually apply the leaf_bgp_asn policy to each applicable switch. When this option is enabled, the system also allows border switches and border gateway switches to use the same ASN without generating a validation error. If this option is later disabled, no validation error is generated for existing switches that already use different ASNs. Validation is enforced only when the leaf_bgp_asn policy is created or updated.

Route-Reflectors (RRs)	The number of spine switches that are used as route reflectors for transporting BGP traffic. Choose 2 or 4 from the drop-down box. The default value is 2. To deploy spine devices as RRs, Nexus Dashboard Fabric Controller sorts the spine devices based on their serial numbers, and designates two or four spine devices as RRs. If you add more spine devices, existing RR configuration won't change. <i>Increasing the count</i> - You can increase the route reflectors from two to four at any point in time. Configurations are automatically generated on the other two spine devices designated as RRs. <i>Decreasing the count</i> - When you reduce four route reflectors to two, remove the unneeded route reflector devices from the fabric. Follow these steps to reduce the count from 4 to 2. 1. Change the value in the drop-down box to 2. 2. Identify the spine switches designated as route reflectors. An instance of the rr_state policy is applied on the spine switch if it's a route reflector. To find out if the policy is applied on the switch, right-click the switch, and choose View/edit policies. In the View/Edit Policies screen, search rr_state in the Template field. It is displayed on the screen. 3. Delete the unneeded spine devices from the fabric (right-click the spine switch icon and choose Discovery > Remove from fabric). If you delete existing RR devices, the next available spine switch is selected as the replacement RR. 4. Click Deploy Config in the fabric topology window. You can preselect RRs and RPs before performing the first Save & Deploy operation. For more information, see <i>Preselecting Switches as Route-Reflectors and Rendezvous-Points</i>.
Anycast Gateway MAC	Specifies the anycast gateway MAC address.
Enable Performance Monitoring	Check the check box to enable performance monitoring. Ensure that you do not clear interface counters from the Command Line Interface of the switches. Clearing interface counters can cause the Performance Monitor to display incorrect data for traffic utilization. If you must clear the counters and the switch has both clear counters and clear counters snmp commands (not all switches have the clear counters snmp command), ensure that you run both the main and the SNMP commands simultaneously. For example, you must run the clear counters interface ethernet slot/port command followed by the clear counters interface ethernet slot/port snmp command. This can lead to a one time spike.

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

completed the necessary configurations for this fabric.

Replication

The fields in the **Replication** tab are described in the following table. Most of the fields are automatically generated based on Cisco-recommended best practice configurations, but you can update the fields if needed.

Field	Description
Replication Mode	The mode of replication that is used in the fabric for BUM (Broadcast, Unknown Unicast, Multicast) traffic. The choices are Ingress Replication or Multicast. When you choose Ingress Replication, the multicast-related fields are disabled. When you choose Multicast replication, the Ingress Replication fields are disabled. You can change the fabric setting from one mode to the other, if no overlay profile exists for the fabric. Choose Multicast as the replication mode for Tenant Routed Multicast (TRM) for IPv4 or IPv6. For more information for the IPv4 use case, see this topic. For more information for the IPv6 use case, see Configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6.
Multicast Subnet	Group IP address prefix used for multicast communication. A unique IP address is allocated from this group for each overlay network. The replication mode change isn't allowed if a policy template instance is created for the current mode. For example, if a multicast-related policy is created and deployed, you can't change the mode to Ingress replication.
IPv6 Multicast Subnet	Enter an IPv6 multicast address with a prefix range of 112 to 126.
Enable Routed (TRM)	IPv4 Tenant Multicast Check this check box to enable Tenant Routed Multicast (TRM) with IPv4 that allows overlay IPv4 multicast traffic to be supported over EVPN/MVPN in VXLAN EVPN fabrics.
Enable Routed (TRM)	IPv6 Tenant Multicast Check this check box to enable Tenant Routed Multicast (TRM) with IPv6 that allows overlay IPv6 multicast traffic to be supported over EVPN/MVPN in VXLAN EVPN fabrics.
Default Address for TRM VRFs	MDT IPv4 The multicast address for Tenant Routed Multicast traffic is populated. By default, this address is from the IP prefix specified in the Multicast Group Subnet field. When you update either field, ensure that the address is chosen from the IP prefix specified in Multicast Group Subnet. For more information, see Overview of Tenant Routed Multicast.

Field	Description
Default MDT IPv6 Address for TRM VRFs	The multicast address for Tenant Routed Multicast traffic is populated. By default, this address is from the IP prefix specified in the IPv6 Multicast Group Subnet field. When you update either field, ensure that you choose the address from the IP prefix specified in IPv6 Multicast Group Subnet. For more information, see Overview of Tenant Routed Multicast.
Rendezvous-Points	Enter the number of spine switches acting as rendezvous points.
RP mode	Choose from the two supported multicast modes of replication, ASM (for Any-Source Multicast [ASM]) or BiDir (for Bidirectional PIM [BIDIR-PIM]). When you choose ASM, the BiDir related fields aren't enabled. When you choose BiDir, the BiDir related fields are enabled.  BIDIR-PIM is supported on Cisco's Cloud Scale Family platforms 9300-EX and 9300-FX/FX2, and software release 9.2(1) onwards. When you create a new VRF for the fabric overlay, this address is populated in the Underlay Multicast Address field, in the Advanced tab.
Underlay RP Loopback ID	The loopback ID used for the rendezvous point (RP), for multicast protocol peering purposes in the fabric underlay.
Underlay Primary RP Loopback ID	Enabled if you choose BIDIR-PIM as the multicast mode of replication. The primary loopback ID used for the phantom RP, for multicast protocol peering purposes in the fabric underlay.
Underlay Backup RP Loopback ID	Enabled if you choose BIDIR-PIM as the multicast mode of replication. The secondary loopback ID used for the phantom RP, for multicast protocol peering purposes in the fabric underlay.
Underlay Second Backup RP Loopback Id	Used for the second fallback Bidir-PIM Phantom RP.
Underlay Third Backup RP Loopback Id	Used for the third fallback Bidir-PIM Phantom RP.
Enable MVPN VRI ID Generation	 This field was previously named Allow L3VNI w/o VLAN under the Advanced tab. In an IPv4 underlay, enabling this option generates an MVPN VRI ID if there is a VRF with TRM or TRMv6 and no VRF with Layer 3 VNI VLAN. In an IPv6 underlay, an MVPN VRI ID is generated once TRM or TRMv6 is enabled, regardless of whether this option is enabled or not. For more information, see Layer 3 VNI without VLAN.

Field	Description
MVPN VRI ID Range	Use this field for allocating a unique MVPN VRI ID per vPC. This field is needed for the following use cases: • If you configure a VXLANv4 underlay with a Layer 3 VNI without VLAN mode, and you enable TRMv4 or TRMv6. • If you configure a VXLANv6 underlay, and you enable TRMv4 or TRMv6.  The MVPN VRI ID cannot be the same as any site ID within a multi-site fabric. The VRI ID has to be unique within all sites within an MSD.
Enable MVPN VRI ID Re-allocation	Enable this check box to generate a one-time VRI ID reallocation. Nexus Dashboard automatically allocates a new MVPN ID within the MVPN VRI ID range above for each applicable switch. Since this is a one-time operation, after performing the operation, this field is turned off.  Changing the VRI ID is disruptive, so plan accordingly.
Auto Configure EVPN Overlay Peering	Check the check box to automatically generate EVPN overlay policies on all leaf and spine switches in the fabric.  When the Auto Configure EVPN Overlay Peering option is enabled, Nexus Dashboard automatically generates the required EVPN overlay policies for all applicable leaf and spine switches in the fabric. If EVPN overlay peering is already configured manually on the switches, customers should manually remove the existing EVPN configurations after enabling this option.

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

vPC

The fields in the **vPC** tab are described in the following table. Most of the fields are automatically generated based on Cisco-recommended best practice configurations, but you can update the fields if needed.

Field	Description
vPC Peer Link VLAN	VLAN used for the vPC peer link SVI.
Make vPC Peer Link VLAN as Native VLAN	Enables vPC peer link VLAN as Native VLAN.

Field	Description
vPC Peer Keep Alive option	Choose the management or loopback option. If you want to use IP addresses assigned to the management port and the management VRF, choose management. If you use IP addresses assigned to loopback interfaces (and a non-management VRF), choose loopback. If you use IPv6 addresses, you must use loopback IDs.
vPC Auto Recovery Time	Specifies the vPC auto recovery time-out period in seconds.
vPC Delay Restore Time	Specifies the vPC delay restore period in seconds.
vPC Peer Link Port Channel ID	Specifies the Port Channel ID for a vPC Peer Link. By default, the value in this field is 500.
vPC IPv6 ND Synchronize	Enables IPv6 Neighbor Discovery synchronization between vPC switches. The check box is enabled by default. Uncheck the check box to disable the function.
vPC advertise-pip	Select the check box to enable the Advertise PIP feature. You can enable the advertise PIP feature on a specific vPC as well.
vPC advertise-pip on Border only	Select the check box to enable advertise-pip on vPC borders and border gateways only. Applicable only when vPC advertise-pip is not enabled.
Enable the same vPC Domain Id for all vPC Pairs	Enable the same vPC Domain ID for all vPC pairs. When you select this field, the vPC Domain Id field is editable.
vPC Domain Id	Specifies the vPC domain ID to be used on all vPC pairs.
vPC Domain Id Range	Specifies the vPC Domain Id range to use for new pairings.
vPC Layer-3 Peer-Router Option	Enable Layer-3 Peer-Router on all leaf switches.
Enable QoS for Fabric vPC-Peering	Enable QoS on spines for guaranteed delivery of vPC fabric peering communication.  QoS for vPC fabric peering and queuing policies options in fabric settings are mutually exclusive.
QoS Policy Name	Specifies QoS policy name that should be same on all fabric vPC peering spines. The default name is spine_qos_for_fabric_vpc_peering .
Use Specific vPC/Port-Channel ID Range	Enable this check box to use a specific vPC/port-channel ID range for leaf-ToR switch pairing.
vPC/Port-Channel ID Range	Specifies one vPC/port-channel ID range for auto-allocating vPC/port-channel IDs for leaf-ToR pairing. The minimum allowed value is 1 and the maximum allowed value is 4096.

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

Protocols

The fields in the **Protocols** tab are described in the following table. Most of the fields are automatically generated based on Cisco-recommended best practice configurations, but you can update the fields if needed.

Field	Description
Underlay Routing Loopback Id	The loopback interface ID is populated as 0 since loopback0 is usually used for fabric underlay IGP peering purposes.
Underlay VTEP Loopback Id	The loopback interface ID is populated as 1 since loopback1 is used for the VTEP peering purposes.
Underlay Anycast Loopback Id	The loopback interface ID is greyed out and used for vPC Peering in VXLANv6 Fabrics only.
Underlay Routing Protocol Tag	The tag defining the type of network.
OSPF Area ID	The OSPF area ID, if OSPF is used as the IGP within the fabric.  The OSPF or IS-IS authentication fields are enabled based on your selection in the Underlay Routing Protocol field in the General tab.
Enable OSPF Authentication	Select the check box to enable OSPF authentication. Deselect the check box to disable it. If you enable this field, the OSPF Authentication Key ID and OSPF Authentication Key fields get enabled.
OSPF Authentication Key ID	The Key ID is populated.
OSPF Authentication Key	The OSPF authentication key must be the 3DES key from the switch.  Plain text passwords are not supported. Log in to the switch, retrieve the encrypted key and enter it in this field. Refer, <i>Retrieving the Authentication Key</i> section for details.
IS-IS Level	Select the IS-IS level from this drop-down list.
Enable IS-IS Network Point-to-Point	Enables network point-to-point on fabric interfaces which are numbered.
Enable IS-IS Authentication	Select the check box to enable IS-IS authentication. Deselect the check box to disable it. If you enable this field, the IS-IS authentication fields are enabled.
IS-IS Authentication Keychain Name	Enter the Keychain name, such as CiscoisisAuth.
IS-IS Authentication Key ID	The Key ID is populated.

Field	Description
IS-IS Authentication Key	Enter the Cisco Type 7 encrypted key.  Plain text passwords are not supported. Log in to the switch, retrieve the encrypted key and enter it in this field. Refer the Retrieving the Authentication Key section for details.
Set IS-IS Overload Bit	When enabled, set the overload bit for an elapsed time after a reload.
IS-IS Overload Bit Elapsed Time	Allows you to clear the overload bit after an elapsed time in seconds.
Enable BGP Authentication	Select the check box to enable BGP authentication. Deselect the check box to disable it. If you enable this field, the BGP Authentication Key Encryption Type and BGP Authentication Key fields are enabled.  If you enable BGP authentication using this field, leave the iBGP Peer-Template Config field blank to avoid duplicate configuration.
BGP Authentication Key Encryption Type	Choose the 3 for 3DES encryption type, or 7 for Cisco encryption type.
BGP Authentication Key	Enter the encrypted key based on the encryption type.  Plain text passwords are not supported. Log in to the switch, retrieve the encrypted key and enter it in the BGP Authentication Key field. Refer the Retrieving the Authentication Key section for details.
Enable PIM Hello Authentication	Select this check box to enable PIM hello authentication on all the intra-fabric interfaces of the switches in a fabric. This check box is editable only for the Multicast replication mode. Note this check box is valid only for the IPv4 underlay.

Field	Description
PIM Hello Authentication Key	Specifies the PIM hello authentication key. For more information, see Retrieving PIM Hello Authentication Key. To retrieve the PIM Hello Authentication Key, perform the following steps: 1. SSH into the switch. 2. On an unused switch interface, enable the following: <pre>switch(config)# interface e1/32 switch(config-if)# ip pim hello-authentication ah-md5 pimHelloPassword</pre> In this example, pimHelloPassword is the cleartext password that has been used. 3. Enter the show run interface command to retrieve the PIM hello authentication key. <pre>switch(config-if)# show run interface e1/32 grep pim ip pim sparse-mode ip pim hello-authentication ah-md5 3 d34e6c5abc7fecf1caa3b588b09078e0</pre> In this example, d34e6c5abc7fecf1caa3b588b09078e0 is the PIM hello authentication key that should be specified in the fabric settings.
Enable BFD	Check the check box to enable feature bfd on all switches in the fabric. This feature is valid only on IPv4 underlay and the scope is within a fabric. BFD within a fabric is supported natively. The BFD feature is disabled by default in the Fabric Settings. If enabled, BFD is enabled for the underlay protocols with the default settings. Any custom required BFD configurations must be deployed via the per switch freeform or per interface freeform policies. The following config is pushed after you select the Enable BFD check box: feature bfd For information about BFD feature compatibility, refer your respective platform documentation and for information about the supported software images, see <i>Compatibility Matrix for Cisco</i>.
Enable BFD for iBGP	Check the check box to enable BFD for the iBGP neighbor. This option is disabled by default.
Enable BFD for OSPF	Check the check box to enable BFD for the OSPF underlay instance. This option is disabled by default, and it is grayed out if the link state protocol is ISIS.

Field	Description
Enable BFD for ISIS	Check the check box to enable BFD for the ISIS underlay instance. This option is disabled by default, and it is grayed out if the link state protocol is OSPF.
Enable BFD for PIM	Check the check box to enable BFD for PIM. This option is disabled by default, and it is be grayed out if the replication mode is Ingress.Following are examples of the BFD global policies: <pre> router ospf <ospf tag> bfd router isis <isis tag> address-family ipv4 unicast bfd ip pim bfd router bgp <bgp asn> neighbor <neighbor ip> bfd </pre>
Enable BFD Authentication	Check the check box to enable BFD authentication. If you enable this field, the BFD Authentication Key ID and BFD Authentication Key fields are editable.  BFD Authentication is not supported when the Fabric Interface Numbering field under the General tab is set to unnumbered. The BFD authentication fields will be grayed out automatically. BFD authentication is valid for only for P2P interfaces.
BFD Authentication Key ID	Specifies the BFD authentication key ID for the interface authentication. The default value is 100.
BFD Authentication Key	Specifies the BFD authentication key. For information about how to retrieve the BFD authentication parameters.

Field	Description
iBGP Peer-Template Config	Add iBGP peer template configurations on the leaf switches to establish an iBGP session between the leaf switch and route reflector. If you use BGP templates, add the authentication configuration within the template and uncheck the Enable BGP Authentication check box to avoid duplicate configuration. In the sample configuration, the 3DES password is displayed after password 3. <pre>router bgp 65000 password 3 sd8478fswerdfw3434fsw4f4w34sdsd8478fswerdfw3434fsw4f4w</pre> The following fields can be used to specify different configurations: • iBGP Peer-Template Config - Specifies the config used for RR and spines with border role. • Leaf/Border/Border Gateway iBGP Peer-Template Config - Specifies the config used for leaf, border, or border gateway. If this field is empty, the peer template defined in iBGP Peer-Template Config is used on all BGP enabled devices (RRs, leafs, border, or border gateway roles). In a brownfield migration, if the spine and leaf use different peer template names, both iBGP Peer-Template Config and Leaf/Border/Border Gateway iBGP Peer-Template Config fields need to be set according to the switch config. If spine and leaf use the same peer template name and content (except for the "route-reflector-client" CLI), only iBGP Peer-Template Config field in fabric setting needs to be set. If the fabric settings on iBGP peer templates do not match the existing switch configuration, an error message is generated and the migration will not proceed.

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

Security

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

Field	Description
Enable Security Groups	Check this check box to enable a security group with an IPv4-only underlay and using the CLI Overlay Mode. For more information on security groups, see Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric.

Field	Description
Security Group Name Prefix	Specify the prefix to use when creating a new security group.
Security Group Tag (SGT) ID Range (Optional)	Specify a tag ID for the security group if necessary.
Security Groups Pre-Provision	Check this check box to generate a security groups configuration for non-enforced VRFs.
Enable MACsec	Check the check box to enable MACsec in the fabric. MACsec configuration is not generated until MACsec is enabled on an intra-fabric link. Perform a Recalculate and deploy operation to generate the MACsec configuration and deploy the configuration on the switch.
MACsec Cipher Suite	Choose one of the following MACsec cipher suites for the MACsec policy: • GCM-AES-128 • GCM-AES-256 • GCM-AES-XPN-128 • GCM-AES-XPN-256 The default value is GCM-AES-XPN-256.
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.  If you disabled the Enable QKD option, you need to specify the MACsec Fallback Key String option.
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
Enable DCI MACsec	Check the check box to enable MACsec on DCI links.  If you enable the Enable DCI MACsec option and disable the Use Link MACsec Setting option on the link, Nexus Dashboard uses the fabric settings for configuring MACsec on the DCI links. For more information on MACsec and using a quantum key distribution (QKD) server, see the section "About connecting two NX-OS fabrics with MACsec using QKD" in Working with Connectivity in Your Nexus Dashboard LAN Fabrics.
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-XPB-128 • GCM-AES-XPB-256 The default value is GCM-AES-XPB-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.

Field	Description
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 the Enable QKD option 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 .
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.

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

Advanced

The fields in the **Advanced** tab are described in the following table. Most of the fields are automatically generated based on Cisco-recommended best practice configurations, but you can update the fields if needed.

Field	Description
VRF Template	Specifies the VRF template for creating VRFs.
Network Template	Specifies the network template for creating networks.
VRF Extension Template	Specifies the VRF extension template for enabling VRF extension to other fabrics.
Network Extension Template	Specifies the network extension template for extending networks to other fabrics.

Field	Description
Overlay Mode	Specify the VRF/network configuration using one of the following options: • config-profile • cli The default option is cli. The overlay mode can only be changed before deploying overlay configurations to the switches. After the overlay configuration is deployed, you cannot change the mode unless all the VRF and network attachments are removed. For a brownfield import, if the overlay is deployed as config-profile, it can be imported in the config-profile mode only. However, if the overlay is deployed as cli, it can be imported in either config-profile or cli modes.
	The Allow L3VNI w/o VLAN option that was previously available under Advanced has been renamed to Enable MVPN VRI ID Generation, and is now available under Replication.
Enable L3VNI w/o VLAN	Check the box to set the default value of the VRF. The setting at this fabric-level field is the default value of Enable L3VNI w/o VLAN at the VRF level. For more information, see: • Layer 3 VNI without VLAN • The "Create a VRF" section in Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric
Enable Private VLAN (PVLAN)	Check this check box to enable private VLAN (PVLAN) on switches, except for spine switches and super spine switches. For more information, see the section "Create private VLANs" section in Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric.
PVLAN Secondary Network Template	Select the template for the PVLAN secondary network. The default is Pvlan_Secondary_Network.
Site ID	The ID for this fabric if you are moving this fabric within an MSD. The site ID is mandatory for a member fabric to be a part of an MSD. Each member fabric of an MSD has a unique site ID for identification.
Intra Fabric Interface MTU	Specifies the MTU for the intra fabric interface. This value should be an even number.
Layer 2 Host Interface MTU	Specifies the MTU for the layer 2 host interface. This value should be an even number.
Unshut Host Interfaces by Default	Check this check box to unshut the host interfaces by default.
Power Supply Mode	Choose the appropriate power supply mode.
CoPP Profile	Choose the appropriate Control Plane Policing (CoPP) profile policy for the fabric. By default, the strict option is populated.

Field	Description
VTEP HoldDown Time	Specifies the NVE source interface hold down time.
Brownfield Overlay Network Name Format	Enter the format to be used to build the overlay network name during a brownfield import or migration. The network name should not contain any white spaces or special characters except underscore (_) and hyphen (-). The network name must not be changed once the brownfield migration has been initiated. For more information, see the section "Create a network" in Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric for the naming convention of the network name. The syntax is [<string> \$\$VLAN_ID\$\$] \$\$VNI\$\$ [<string> \$\$VLAN_ID\$\$] and the default value is Auto_Net_VNI\$\$VNI\$\$_VLAN\$\$VLAN_ID\$\$. When you create networks, the name is generated according to the syntax you specify. The following list describes the variables in the syntax: ▪ \$\$VNI\$\$: Specifies the network VNI ID found in the switch configuration. This is a mandatory keyword required to create unique network names. ▪ \$\$VLAN_ID\$\$: Specifies the VLAN ID associated with the network. VLAN ID is specific to switches, hence Nexus Dashboard Fabric Controller picks the VLAN ID from one of the switches, where the network is found, randomly and use it in the name. We recommend not to use this unless the VLAN ID is consistent across the fabric for the VNI. ▪ <string>: This variable is optional and you can enter any number of alphanumeric characters that meet the network name guidelines. An example overlay network name: Site_VNI12345_VLAN1234  Ignore this field for greenfield deployments. The Brownfield Overlay Network Name Format applies for the following brownfield imports: ▪ CLI-based overlays ▪ Configuration profile-based overlay
Skip Overlay Network Interface Attachments	Check the check box to skip overlay network interface attachments for Brownfield and Host Port Resync cases.
Enable CDP for Bootstrapped Switch	Enables CDP on management (mgmt0) interface for bootstrapped switch. By default, for bootstrapped switches, CDP is disabled on the mgmt0 interface.

Field	Description
Enable VXLAN OAM	Enables the VXLAN OAM functionality for devices in the fabric. This is enabled by default. Uncheck the check box to disable VXLAN OAM function. If you want to enable the VXLAN OAM function on specific switches and disable on other switches in the fabric, you can use freeform configurations to enable OAM and disable OAM in the fabric settings.  The VXLAN OAM feature in Cisco Nexus Dashboard Fabric Controller is only supported on a single fabric or site.
Enable Tenant DHCP	Check the check box to enable feature dhcp and associated configurations globally on all switches in the fabric. This is a pre-requisite for support of DHCP for overlay networks that are part of the tenant VRFs.  Ensure that Enable Tenant DHCP is enabled before enabling DHCP-related parameters in the overlay profiles.
Enable NX-API	Specifies enabling of NX-API on HTTPS. This check box is checked by default.
NX-API HTTPS Port Number	Field becomes active if the Enable NX-API option is enabled. Enter the NX-API HTTPS port number. Default value is 443.
Enable NX-API on HTTP Port	Specifies enabling of NX-API on HTTP. Enable this check box and the Enable NX-API check box to use HTTP. This check box is checked by default. If you uncheck this check box, the applications that use NX-API and supported by Cisco Nexus Dashboard Fabric Controller, such as Endpoint Locator (EPL), Layer 4-Layer 7 services (L4-L7 services), VXLAN OAM, and so on, start using the HTTPS instead of HTTP.  If you check the Enable NX-API check box and the Enable NX-API on HTTP check box, applications use HTTP.
NX-API HTTP Port Number	Field becomes active if the Enable HTTP NX-API option is enabled. Enter the NX-API HTTPS port number. Default value is 80.
Enable L4-L7 Services Re-direction	Check this check box to enable the following features on the switch, based on the L4-L7 services use case: • Policy-based redirect (PBR) • Enhanced policy-based redirect (ePBR) • Service level agreement (SLA) sender
Enable Strict Config Compliance	Enable the Strict Config Compliance feature by selecting this check box. It enables bi-directional compliance checks to flag additional configs in the running config that are not in the intent/expected config. By default, this feature is disabled.

Field	Description
Enable AAA IP Authorization	Enables AAA IP authorization, when IP Authorization is enabled in the remote authentication server. This is required to support Nexus Dashboard Fabric Controller in scenarios where customers have strict control of which IP addresses can have access to the switches.
Enable NDFC as Trap Host	Select this check box to enable Nexus Dashboard Fabric Controller as an SNMP trap destination. Typically, for a native HA Nexus Dashboard Fabric Controller deployment, the eth1 VIP IP address will be configured as SNMP trap destination on the switches. By default, this check box is enabled.
Anycast Border Gateway advertise-pip	Enables to advertise Anycast Border Gateway PIP as VTEP. Effective on MSD fabric 'Recalculate Config'.
Greenfield Cleanup Option	Enable the switch cleanup option for switches imported into Nexus Dashboard Fabric Controller with Preserve-Config=No, without a switch reload. This option is typically recommended only for the fabric environments with Cisco N9000v Switches to improve on the switch clean up time. The recommended option for Greenfield deployment is to employ Bootstrap or switch cleanup with a reboot. In other words, this option should be unchecked.
Enable Precision Time Protocol (PTP)	Enables PTP across a fabric. When you check this check box, PTP is enabled globally and on core-facing interfaces. Additionally, the PTP Source Loopback Id and PTP Domain Id fields are editable. For more information, see the section "Precision Time Protocol for Data Center VXLAN EVPN Fabrics" in Precision Time Protocol for Data Center VXLAN EVPN fabrics .
PTP Source Loopback Id	Specifies the loopback interface ID Loopback that is used as the Source IP Address for all PTP packets. The valid values range from 0 to 1023. The PTP loopback ID cannot be the same as RP, Phantom RP, NVE, or MPLS loopback ID. Otherwise, an error will be generated. The PTP loopback ID can be the same as BGP loopback or user-defined loopback which is created from Nexus Dashboard Fabric Controller. If the PTP loopback ID is not found during Deploy Config, the following error is generated: Loopback interface to use for PTP source IP is not found. Create PTP loopback interface on all the devices to enable PTP feature.
PTP Domain Id	Specifies the PTP domain ID on a single network. The valid values range from 0 to 127.
PTP Source VLAN Id	Specifies the SVI used for PTP source on ToRs. The valid values range from 2 to 3967.
Enable MPLS Handoff	Check the check box to enable the MPLS Handoff feature. For more information, see Overview of VXLAN EVPN to SR-MPLS and MPLS LDP interconnection .
Underlay MPLS Loopback Id	Specifies the underlay MPLS loopback ID. The default value is 101.

Field	Description
IS-IS NET Area Number for MPLS Handoff	Specify the IS-IS NET area number for the MPLS handoff.
Enable TCAM Allocation	TCAM commands are automatically generated for VXLAN and vPC Fabric Peering when enabled.
Enable Default Queuing Policies	Check this check box to apply QoS policies on all the switches in this fabric. To remove the QoS policies that you applied on all the switches, uncheck this check box, update all the configurations to remove the references to the policies, and save and deploy. Pre-defined QoS configurations are included that can be used for various Cisco N9000 Series Switches. When you check this check box, the appropriate QoS configurations are pushed to the switches in the fabric. The system queuing is updated when configurations are deployed to the switches. You can perform the interface marking with defined queuing policies, if required, by adding the required configuration to the per interface freeform block. Review the actual queuing policies by opening the policy file in the template editor. From Cisco Nexus Dashboard Fabric Controller Web UI, choose Manage > Templates. Search for the queuing policies by the policy file name, for example, queuing_policy_default_8q_cloudscale. Choose the file. From the Actions drop-down list, select Edit template content to edit the policy. For more information, see the <i>Cisco N9000 Series NX-OS Quality of Service Configuration Guide</i> for platform specific details.
N9K Cloud Scale Platform Queuing Policy	Choose the queuing policy from the drop-down list to be applied to all Cisco N9200 Series Switches and the Cisco N9000 Series Switches that end with EX, FX, and FX2 in the fabric. The valid values are queuing_policy_default_4q_cloudscale and queuing_policy_default_8q_cloudscale . Use the queuing_policy_default_4q_cloudscale policy for FEXes. You can change from the queuing_policy_default_4q_cloudscale policy to the queuing_policy_default_8q_cloudscale policy only when FEXes are offline.
N9K R-Series Platform Queuing Policy	Choose the queuing policy from the drop-down list to be applied to all Cisco Nexus switches that ends with R in the fabric. The valid value is queuing_policy_default_r_series .
Other N9K Platform Queuing Policy	Choose the queuing policy from the drop-down list to be applied to all other switches in the fabric other than the switches mentioned in the above two options. The valid value is queuing_policy_default_other .

Field	Description
Priority flow control watch-dog interval	When enabling the AI feature, priority-flow-control watch-dog-interval on is enabled on all of your configured devices, intra-fabric links, and all your host interfaces where PFC is also enabled. This release also adds the Priority flow control watch-dog interval field. Here you can set the Priority flow control watch-dog interval field to a non-system default value (default is 100 milliseconds). Valid values are <101-1000>.  The watch-dog is not supported on Cisco N9K-C9808 and Cisco N9K-C9804 series switches on NXOS version 10.5(1).
Enable Real Time Interface Statistics Collection	Valid for NX-OS fabrics only. Check the box to enable the collection of real time interface statistics.
Interface Statistics Load Interval	Enter the interval for the interface statistics load, in seconds (min: 5, max: 300).
Spanning Tree Root Bridge Protocol	Specify the protocol to be used for configuring Root Bridge: Options are: • rpvst+: Rapid Per-VLAN Spanning Tree • mst: Multiple Spanning Tree • unmanaged (default): STP Root not managed by Nexus Dashboard  Spanning Tree settings and bridge configurations are applicable at the Aggregation layer only.
Spanning Tree VLAN Range	Specify the VLAN range. For example: 1, 3-5, 7, 9-11 The default value is 1-3967. Applicable only for Aggregation devices.
MST Instance Range	Specify the MST instance range. For example: 0-3,5,7-9 The default value is 0. Applicable only for Aggregation devices.
Spanning Tree Bridge Priority	Specify the bridge priority for the spanning tree in increments of 4096. Applicable only for Aggregation devices.

Field	Description
Set Allowed Vlan On Leaf-ToR Pairing	The new fabric parameter Set Allowed Vlan On Leaf-ToR Pairing sets the trunk allowed VLAN to none or all on all leaf switch-to-ToR pairing port-channels. With this new fabric parameter, a manual change of the Trunk Allowed Vlans parameter on a leaf switch-to-ToR pairing port-channel is no longer supported.  If you manually modified the allowed VLANs to all on the uplink_access port-channels on the ToRs on a release prior to Nexus Dashboard release 4.1.1, use leaf switch-to-ToR pairing instead of uplink_access. This issue can happen on leaf switches or ToRs.

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

Freeform

The fields in this tab are shown below. For more information, see "Enable freeform configurations on fabric switches" in [Working with Inventory in Your Nexus Dashboard LAN or IPFM Fabrics](#).

Field	Description
Leaf Pre-Interfaces Freeform Config	Enter additional CLIs, added before interface configurations, for all Leafs and Tier2 Leafs as captured from Show Running Configuration.
Spine Pre-Interfaces Freeform Config	Enter additional CLIs, added before interface configurations, for all Spines as captured from Show Running Configuration.
ToR Pre-Interfaces Freeform Config	Enter additional CLIs, added before interface configurations, for all ToRs as captured from Show Running Configuration.
Leaf Post-Interfaces Freeform Config	Enter additional CLIs, added after interface configurations, for all Leafs and Tier2 Leafs as captured from Show Running Configuration.
Spine Post-Interfaces Freeform Config	Enter additional CLIs, added after interface configurations, for all Spines as captured from Show Running Configuration.
ToR Post-Interfaces Freeform Config	Enter additional CLIs, added after interface configurations, for all ToRs as captured from Show Running Configuration.
Intra-fabric Links Additional Config	Add CLIs that should be added to the intra-fabric links.

Resources

The fields in the **Resources** tab are described in the following table. Most of the fields are automatically generated based on Cisco-recommended best practice configurations, but you can update the fields if needed.

Field	Description
Manual Underlay IP Address Allocation	<i>Do not</i> check this check box if you are transitioning your VXLAN fabric management to Nexus Dashboard Fabric Controller. - By default, Nexus Dashboard Fabric Controller allocates the underlay IP address resources (for loopbacks, fabric interfaces, etc) dynamically from the defined pools. If you check the check box, the allocation scheme switches to static, and some of the dynamic IP address range fields are disabled. - For static allocation, the underlay IP address resources must be populated into the Resource Manager (RM) using REST APIs. - The Underlay RP Loopback IP Range field stays enabled if BIDIR-PIM function is chosen for multicast replication. - Changing from static to dynamic allocation keeps the current IP resource usage intact. Only future IP address allocation requests are taken from dynamic pools.
Underlay Routing Loopback IP Range	Specifies loopback IP addresses for the protocol peering.
Underlay VTEP Loopback IP Range	Specifies loopback IP addresses for VTEPs.
Underlay RP Loopback IP Range	Specifies the anycast or phantom RP IP address range.
Underlay Subnet IP Range	IP addresses for underlay P2P routing traffic between interfaces.
Underlay MPLS Loopback IP Range	Specifies the underlay MPLS loopback IP address range. For eBGP between Border of Easy A and Easy B, Underlay routing loopback and Underlay MPLS loopback IP range must be a unique range. It should not overlap with IP ranges of the other fabrics, else VPNv4 peering will not come up.
Underlay Routing Loopback IPv6 Range	Specifies Loopback0 IPv6 Address Range
Underlay VTEP Loopback IPv6 Range	Specifies Loopback1 and Anycast Loopback IPv6 Address Range.
Underlay Subnet IPv6 Range	Specifies IPv6 Address range to assign Numbered and Peer Link SVI IPs.
BGP Router ID Range for IPv6 Underlay	Specifies BGP router ID range for IPv6 underlay.
Layer 2 VXLAN VNI Range	Specifies the overlay VXLAN VNI range for the fabric (min:1, max:16777214).
Layer 3 VXLAN VNI Range	Specifies the overlay VRF VNI range for the fabric (min:1, max:16777214).
Network VLAN Range	VLAN range for the per switch overlay network (min:2, max:4094).
VRF VLAN Range	VLAN range for the per switch overlay Layer 3 VRF (min:2, max:4094).

Field	Description
Subinterface Range Dot1q	Specifies the subinterface range when L3 sub interfaces are used.
VRF Lite Deployment	Specify the VRF Lite method for extending inter fabric connections. The VRF Lite Subnet IP Range field specifies resources reserved for IP address used for VRF Lite when VRF Lite IFCs are auto-created. If you select Back2Back&ToExternal, then VRF Lite IFCs are auto-created.
Auto Deploy for Peer	This check box is applicable for VRF Lite deployment. When you select this checkbox, auto-created VRF Lite IFCs will have the Auto Generate Configuration for Peer field in the VRF Lite tab set. To access VRF Lite IFC configuration, navigate to the Links tab, select the particular link, and then choose Actions > Edit. You can check or uncheck the check box when the VRF Lite Deployment field is not set to Manual. This configuration only affects the new auto-created IFCs and does not affect the existing IFCs. You can edit an auto-created IFC and check or uncheck the Auto Generate Configuration for Peer field. This setting takes priority always.
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 check box 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, 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 select this check box, 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 check box 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.  To access the Auto Generate Configuration on default VRF and Auto Generate Configuration for NX-OS Peer on default VRF fields for an IFC link, navigate to the Links tab, select the particular link and choose Actions > Edit.
Redistribute Route-map Name BGP	Defines the route map for redistributing the BGP routes in default VRF.

Field	Description
VRF Lite Subnet IP Range and VRF Lite Subnet Mask	These fields are prefilled with the DCI subnet details. Update the fields as needed. The values shown on the page are automatically generated. If you want to update the IP address ranges, VXLAN Layer 2/Layer 3 network ID ranges or the VRF/network VLAN ranges, ensure that each fabric has its own unique range and is distinct from any underlay range to avoid possible duplication. You should only update one range of values at a time. If you want to update more than one range of values, do it in separate instances. For example, if you want to update Layer 2 and Layer 3 ranges, you should do the following. 1. Update the Layer 2 range and click Save. 2. Click the Edit Fabric option again, update the Layer 3 range, and click Save.
Service Network VLAN Range	Specifies a VLAN range in the Service Network VLAN Range field. This is a per switch overlay service network VLAN range. The minimum allowed value is 2 and the maximum allowed value is 3967.
Auto Allocation of Unique IP on VRF Extension over VRF Lite IFC	Automatically allocates a unique IPv4 address with subnet for the source and the destination interfaces for VRF extensions over VRF Lite IFC. When enabled, the system auto populates a unique IP address for the source and the destination interfaces for each extension in the VRF attachment. When you disable the feature, the system auto populates the same IP address for the source and the destination interfaces for the VRF extensions and these IP addresses are allocated in resource manager with the VRFs attached. The resource manager ensures that they are not used for any other purpose on the same VRF.

Field	Description
Per VRF Per VTEP Loopback IPv4 Auto-Provisioning	Auto provisions a loopback address in IPv4 format for a VTEP that the system uses for VRF attachment. This option is not enabled by default. When enabled, the system allocates an IPv4 address from the IP pool that you have assigned for the VTEP loopback interface. 1. Enable the Per VRF Per VTEP Loopback IPv4 Auto-Provisioning or the Per VRF Per VTEP Loopback IPv6 Auto-Provisioning option. 2. Save the fabric settings. 3. Perform a Recalculate and Deploy operation. 4. Navigate to VRF Attachments. 5. If certain VRFs are already attached, click Actions > Quick attach. This generates the new loopback in the VRF.  If VRF extensions are already enabled and configured, for example, VRF Lite on a border device, prior to enabling the fabric setting, you need to access the respective VRF attachment and the border device to reattach the VRF extension again. For example, VRF Corp is attached on Border-1 and extended to an external domain using VRF Lite. In this situation, when you perform a Quick attach to provision the new loopback in the VRF, the original VRF-Lite extension gets detached. You can then select the VRF attachment, edit, and re-attach the VRF-Lite extension and then deploy all the relevant configurations.
Per VRF Per VTEP IPv4 Pool for Loopbacks	A pool of IPv4 addresses assigned to the loopback interfaces on VTEPs for each VRF.

Field	Description
Per VRF Per VTEP Unique Loopback IPv4 Auto-Provisioning	Auto provisions a unique loopback address in IPv4 format for a VTEP that the system uses for VRF attachment. This option is not enabled by default. When enabled, the system allocates an IPv4 address from the IP pool that you have assigned for the VTEP loopback interface. The allocated loopback IP address is unique across all VRFs. 1. Enable the Per VRF Per VTEP Unique Loopback IPv4 Auto-Provisioning or the Per VRF Per VTEP Unique Loopback IPv6 Auto-Provisioning option. 2. Save the fabric settings. 3. Perform a Recalculate and Deploy operation. 4. Navigate to VRF Attachments. 5. If certain VRFs are already attached, click Actions > Quick attach. This generates the new loopback in the VRF. If VRF extensions are already enabled and configured, for example, VRF Lite on a border device, prior to enabling the fabric setting, you must access the respective VRF attachment and the border device to reattach the VRF extension. For example, if VRF Corp is attached on Border-1 and extended to an external domain by using VRF Lite, when you perform a Quick attach operation to provision the new loopback in the VRF, the original VRF-Lite extension is detached. You can then select the VRF attachment, edit it, reattach the VRF-Lite extension, and deploy the relevant configurations.  The Per VRF Per VTEP Loopback Auto-Provisioning fields and the Per VRF Per VTEP Unique Loopback Auto-Provisioning fields are mutually exclusive. Use only one loopback auto-provisioning method for a fabric.
Per VRF Per VTEP IPv4 Pool for Unique Loopbacks	A pool of IPv4 addresses assigned to the loopback interfaces on VTEPs for each VRF. Each VRF uses unique loopback addresses from this pool for the applicable VTEPs.
Per VRF Per VTEP Unique Loopback IPv6 Auto-Provisioning	Auto provisions a unique loopback address in IPv6 format for a VTEP that the system uses for VRF attachment. This option is not enabled by default. When enabled, the system allocates an IPv6 address from the IP pool that you have assigned for the VTEP loopback interface.
Per VRF Per VTEP IPv6 Pool for Unique Loopbacks	A pool of IPv6 addresses assigned to the loopback interfaces on VTEPs for each VRF.

Field	Description
Service Level Agreement (SLA) ID Range	The per switch SLA ID Range (Min:1, Max: 2147483647).
Tracked Object ID Range	The per switch tracked object ID Range (Min:1, Max: 512).
Service Network VLAN Range	The per switch Overlay Service Network VLAN Range (Min:2, Max:4094).
Route Map Sequence Number Range	Specifies the route map sequence number range. The minimum allowed value is 1 and the maximum allowed value is 65534.

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

Manageability

The fields in the **Manageability** tab are described in the following table. Most of the fields are automatically generated based on Cisco-recommended best practice configurations, but you can update the fields if needed.

Field	Description
Inband Management	Enabling this allows the management of the switches over their front panel interfaces. The Underlay Routing Loopback interface is used for discovery. If enabled, switches cannot be added to the fabric over their out-of-band (OOB) mgmt0 interface. To manage easy fabrics through Inband management, ensure that you have chosen Data in Nexus Dashboard Web UI, Admin > System Settings > Server Settings > Admin. While both inband management and out-of-band connectivity (mgmt0) are supported for this setting, inband management is not supported when discovering or onboarding brownfield Data Center VXLAN EVPN fabrics. You can switch to inband management after discovering or onboarding brownfield Data Center VXLAN EVPN fabrics, but onboarding must happen over out-of-band. For more information, see Configuring Inband Management and Out-of-Band PnP.
DNS Server IPs	Specifies the comma separated list of IP addresses (v4/v6) of the DNS servers.
DNS Server VRFs	Specifies one VRF for all DNS servers or a comma separated list of VRFs, one per DNS server.
NTP Server IPs/Hostnames	Specifies a comma-separated list of IP addresses (IPv4/IPv6) or hostnames for the NTP server. Hostnames are limited to 80 characters in length and must not contain any whitespace or special characters, except for hyphens (-) and periods (.).
NTP Server VRFs	Specifies one VRF for all NTP servers or a comma separated list of VRFs, one per NTP server.

Field	Description
Syslog Server IPs/Hostnames	Specifies a comma-separated list of IP addresses (IPv4/IPv6) or hostnames for the Syslog server. Hostnames are limited to 199 characters in length and should not contain any whitespace or special characters, except for hyphens (-) and periods (.).
Syslog Server Severity	Specifies the comma separated list of syslog severity values, one per syslog server. The minimum value is 0 and the maximum value is 7. To specify a higher severity, enter a higher number.
Syslog Server VRFs	Specifies one VRF for all syslog servers or a comma separated list of VRFs, one per syslog server.
AAA Freeform Config	Specifies the AAA freeform configurations. If AAA configurations are specified in the fabric settings, switch_freeform PTI with source as UNDERLAY_AAA and description as AAA Configurations will be created.
Banner	Use the Banner field to set a login banner that Nexus Dashboard applies to all switches in the fabric. This banner usually includes legal warnings, security notices, and compliance information. If you enter non-banner CLI commands in this field, Nexus Dashboard will ignore them or cause deployment to fail with a validation error.  Manually update the last line of the Banner configuration so that the # and the ! are on the same line: #!. If the # and the ! remain on separate lines, it prevents configurations from deploying to those devices.

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

Bootstrap

The fields in the **Bootstrap** tab are described in the following table. Most of the fields are automatically generated based on Cisco-recommended best practice configurations, but you can update the fields if needed.

Field	Description
Enable Bootstrap	Select this check box to enable the bootstrap feature. Bootstrap allows easy day-0 import and bring-up of new devices into an existing fabric. Bootstrap leverages the NX-OS POAP functionality. To add more switches and for POAP capability, choose the check boxes for Enable Bootstrap and Enable Local DHCP Server. For more information, see Configuring Inband Management and Out-of-Band PnP. After you enable bootstrap, you can enable the DHCP server for automatic IP address assignment using one of the following methods: • External DHCP Server: Enter information about the external DHCP server in the Switch Mgmt Default Gateway and Switch Mgmt IP Subnet Prefix fields. • Local DHCP Server: Enable the Local DHCP Server check box and enter details for the remaining mandatory fields.
Enable Local DHCP Server	Select this check box to initiate enabling of automatic IP address assignment through the local DHCP server. When you select this check box, the DHCP Scope Start Address and DHCP Scope End Address fields become editable. If you do not select this check box, Nexus Dashboard Fabric Controller uses the remote or external DHCP server for automatic IP address assignment.
DHCP Version	Select DHCPv4 or DHCPv6 from this drop-down list. When you select DHCPv4, the Switch Mgmt IPv6 Subnet Prefix field is disabled. If you select DHCPv6, the Switch Mgmt IP Subnet Prefix is disabled.  Cisco N9000 and Cisco Nexus 3000 Series Switches support IPv6 POAP only when switches are either Layer-2 adjacent (eth1 or out-of-band subnet must be a /64) or they are L3 adjacent residing in some IPv6 /64 subnet. Subnet prefixes other than /64 are not supported.
DHCP Scope Start Address and DHCP Scope End Address	Specifies the first and last IP addresses of the IP address range to be used for the switch out of band POAP.
Switch Mgmt Default Gateway	Specifies the default gateway for the management VRF on the switch.
Switch Mgmt IP Subnet Prefix	Specifies the prefix for the Mgmt0 interface on the switch. The prefix should be between 8 and 30. <i>DHCP scope and management default gateway IP address specification -</i> If you specify the management default gateway IP address 10.0.1.1 and subnet mask 24, ensure that the DHCP scope is within the specified subnet, between 10.0.1.2 and 10.0.1.254.

Field	Description
Switch Mgmt IPv6 Subnet Prefix	Specifies the IPv6 prefix for the Mgmt0 interface on the switch. The prefix should be between 112 and 126. This field is editable if you enable IPv6 for DHCP.
Enable AAA Config	Select this check box to include AAA configurations from the Manageability tab as part of the device start-up config post bootstrap.
DHCPv4/DHCPv6 Multi Subnet Scope	Specifies the field to enter one subnet scope per line. This field is editable after you check the Enable Local DHCP Server check box. The format of the scope should be defined as: DHCP Scope Start Address, DHCP Scope End Address, Switch Management Default Gateway, Switch Management Subnet Prefix For example: 10.6.0.2, 10.6.0.9, 10.6.0.1, 24
Bootstrap Config Freeform	(Optional) Enter additional commands as needed. For example, if you require some additional configurations to be pushed to the device and be available post device bootstrap, they can be captured in this field, to save the desired intent. After the devices boot up, they will contain the configuration defined in the Bootstrap Freeform Config field. Copy-paste the running-config to a freeform config field with correct indentation, as seen in the running configuration on the NX-OS switches. The freeform config must match the running config. For more information, see the section "Enable freeform configurations on fabric switches" in Working with Inventory in Your Nexus Dashboard LAN or IPFM Fabrics.

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

Configuration Backup

The fields in the **Configuration Backup** tab are described in the following table. Most of the fields are automatically generated based on Cisco-recommended best practice configurations, but you can update the fields if needed.

Field	Description
Hourly Fabric Backup	Select the check box to enable an hourly backup of fabric configurations and the intent. The hourly backups are triggered during the first 10 minutes of the hour.
Scheduled Backup Fabric	Check the check box to enable a daily backup. This backup tracks changes in running configurations on the fabric devices that are not tracked by configuration compliance.

Field	Description
Scheduled Time	Specify the scheduled backup time in a 24-hour format. This field is enabled if you check the Scheduled Fabric Backup check box. Select both the check boxes to enable both back up processes. The backup process is initiated after you click Save. The scheduled backups are triggered exactly at the time you specify with a delay of up to two minutes. The scheduled backups are triggered regardless of the configuration deployment status. The number of fabric backups that will be retained on Nexus Dashboard is decided by the Admin > System Settings > Server Settings > LAN Fabric > Maximum Backups per Fabric. The number of archived files that can be retained is set in the # Number of archived files per device to be retained: field in the Server Properties window.  To trigger an immediate backup, do the following: 1. Choose Overview > Topology. 2. Click within the specific fabric box. The fabric topology screen comes up. 3. Right-click on a switch within the fabric, then select Preview Config. 4. In the Preview Config window for this fabric, click Re-Sync All. You can also initiate the fabric backup in the fabric topology window. Click Backup Now in the Actions pane.

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

Flow Monitor

The fields in the **Flow Monitor** tab are described in the following table. Most of the fields are automatically generated based on Cisco-recommended best practice configurations, but you can update the fields if needed.

Field	Description
Enable Netflow	Check this check box to enable Netflow on VTEPs for this fabric. By default, Netflow is disabled. When enabled, NetFlow configuration will be applied to all VTEPS that support netflow.  When Netflow is enabled on the fabric, you can choose not to have netflow on a particular switch by having a dummy no_netflow PTI. If netflow is not enabled at the fabric level, an error message is generated when you enable netflow at the interface, network, or vrf level. For information about Netflow support for Cisco Nexus Dashboard, see the "Configuring Netflow support" section in Creating Fabrics and Fabric Groups.

In the **Netflow Exporter** area, choose **Actions > Add** to add one or more Netflow exporters. This exporter is the receiver of the netflow data. The fields on this page are:

- **Exporter Name** - Specifies the name of the exporter.
- **IP** - Specifies the IP address of the exporter.
- **VRF** - Specifies the VRF over which the exporter is routed.
- **Source Interface** - Specifies the source interface name.
- **UDP Port** - Specifies the UDP port over which the netflow data is exported.

Click **Save** to configure the exporter. Click **Cancel** to discard. You can also choose an existing exporter and choose **Actions > Edit** or **Actions > Delete** to perform relevant actions.

In the **Netflow Record** area, click **Actions > Add** to add one or more Netflow records. The fields on this screen are:

- **Record Name** - Specifies the name of the record.
- **Record Template** - Specifies the template for the record. Enter one of the record templates names. In Release 12.0.2, the following two record templates are available for use. You can create custom netflow record templates. Custom record templates saved in the template library are available for use here.
 - **netflow_ipv4_record** - Uses the IPv4 record template.
 - **netflow_l2_record** - Uses the Layer 2 record template.
- **Is Layer2 Record** - Check this check box if the record is for Layer2 netflow.

Click **Save** to configure the report. Click **Cancel** to discard. You can also choose an existing record and select **Actions > Edit** or **Actions > Delete** to perform relevant actions.

In the **Netflow Monitor** area, click **Actions > Add** to add one or more Netflow monitors. The fields on this page are:

- **Monitor Name** - Specifies the name of the monitor.
- **Record Name** - Specifies the name of the record for the monitor.

- **Exporter1 Name** - Specifies the name of the exporter for the netflow monitor.
- **Exporter2 Name** - (optional) Specifies the name of the secondary exporter for the netflow monitor.

The record name and exporters referred to in each netflow monitor must be defined in "**Netflow Record**" and "**Netflow Exporter**".

Click **Save** to configure the monitor. Click **Cancel** to discard. You can also choose an existing monitor and select **Actions > Edit** or **Actions > Delete** to perform relevant actions.

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

Telemetry

For a Nexus Dashboard fabric in Monitored mode, Nexus Dashboard will not deploy the telemetry configuration to the switches in the fabric. In order to activate, pause, or deactivate telemetry, you must copy and paste the expected configuration on the switches.



1. Navigate to **Admin > System Status > Telemetry > Switches**
2. Click the checkbox next to a switch, then click **Actions > Expected configuration**.
3. View and copy the configuration information in the **Software Telemetry** and **Flow Telemetry** area,.
4. Using the command line, log in to the switch.
5. Enter this command.

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

The telemetry feature in Nexus Dashboard allows you to collect, manage, and monitor real-time telemetry data from your Nexus Dashboard. This data provides valuable insights into the performance and health of your network infrastructure, enabling you to troubleshoot proactively and optimize operations. When you enable telemetry, you gain enhanced visibility into network operations and efficiently manage your fabrics.

Follow these steps to enable telemetry for a specific fabric.

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

Go to **Manage > Fabrics**.

2. Choose the fabric for which you want to enable telemetry.
3. From the **Actions** drop-down list, choose **Edit fabric settings**.

The **Edit *fabric-name* settings** page displays.



You can also access the **Edit *fabric-name* settings** page for a fabric from the **Fabric Overview** page. In the **Fabric Overview** page, click the **Actions** drop-down list and choose **Edit fabric settings**.

4. In the **Edit *fabric-name* settings** page, click the **General** tab.
5. Under the **Enabled features** section, check the **Telemetry** check box.
6. Click **Save**.

Navigate back to the **Edit *fabric-name* settings** page. The **Telemetry** tab displays.



The **Telemetry** tab appears only when you enable the **Telemetry** option under the **General** tab in the **Edit *fabric-name* settings** page.

The **Telemetry** tab includes these options.

- **Configuration**—allows you to manage telemetry settings and parameters.
- **NAS**—provides Network Analytics Service (NAS) features for advanced insights.

Edit configuration settings

The **Configuration** tab includes these settings.

- **General**—allows you to enable analysis.

You can enable these settings.

- **Enable assurance analysis**—enables you to collect of telemetry data from devices to ensure network reliability and performance.
- **Enable Microburst sensitivity** - allows you to monitor traffic to detect unexpected data bursts within a very small time window (microseconds). Choose the sensitivity type from the **Microburst Sensitivity Level** drop-down list. The options are **High sensitivity**, **Medium sensitivity**, and **Low sensitivity**.



The **Enable Microburst sensitivity** option is available only for ACI fabrics.

- **Flow collection modes**—allows you to choose the mode for telemetry data collection. Modes include **NetFlow**, **sFlow**, and **Flow Telemetry**.

For more information see: [Flow collection](#) and [Configure flows](#).

- **Flow collection rules**—allows you to define rules for monitoring specific subnets or endpoints. These rules are pushed to the relevant devices, enabling detailed telemetry data collection.

For more information, see [Flow collection](#).

Edit NAS settings

Nexus Dashboard allows you to export captured flow records to a remote NAS device using the Network File System (NFS) protocol. Nexus Dashboard defines the directory structure on NAS where the flow records are exported.

You can choose between the two export modes.

- **Full**—exports the complete data for each flow record.
- **Base**—exports only the essential 5-tuple data for each flow record.

Nexus Dashboard needs both read and write permissions on the NAS to perform the export successfully. If Nexus Dashboard cannot write to the NAS, it will generate an alert to notify you of the issue.

Disable Telemetry

You can uncheck the **Telemetry** check box on your fabric's **Edit Fabric Settings > General >** page to disable the telemetry feature for your fabric. Disabling telemetry puts the telemetry feature in a transition phase and eventually the telemetry feature is disabled.

In certain situations, the disable telemetry workflow can fail, and you may see the **Force disable telemetry** option on your fabric's **Edit Fabric Settings** page.

If you disable the telemetry option using the instructions provided in [Perform force disable telemetry on your fabric](#), Nexus Dashboard acknowledges the user intent to disable telemetry feature for your fabric, ignoring any failures.

The Nexus Dashboard **Force disable telemetry** allows you to perform a force disable action for the telemetry configuration on your fabric. This action is recommended when the telemetry disable workflow has failed and you need to disable the telemetry feature on your fabric.



Using the **Force disable telemetry** feature may leave switches in your fabric with stale telemetry configurations. You must manually clean up these stale configurations on the switches before re-enabling telemetry on your fabric.

Perform force disable telemetry on your fabric

Follow these steps to perform a force disable telemetry on your fabric.

1. (Optional) Before triggering a force disable of telemetry configuration, resolve any telemetry configuration anomalies flagged on the fabric.
2. On the **Edit Fabric Settings** page of your fabric, a banner appears to alert you that telemetry cannot be disabled gracefully, and a **Force Disable** option is provided with the alert message.
3. Disable telemetry from the Nexus Dashboard UI using one of these options.
 - a. Click the **Force disable** option in the banner that appears at the top of your fabric's **Edit Fabric Settings** page to disable telemetry for your fabric gracefully.
 - b. Navigate to your fabric's **Overview** page and click the **Actions** drop-down list to choose **Telemetry > Force disable telemetry** option.

Once the force disable action is executed, the **Telemetry** configuration appears as disabled in **Edit Fabric Settings > General > Enabled features > Telemetry** area, that is, the **Telemetry** check box is unchecked.

4. Clean up any stale telemetry configurations from the fabric before re-enabling telemetry on Nexus Dashboard.

NAS

You can export flow records captured by Nexus Dashboard on a remote Network Attached Storage (NAS) with NFS.

Nexus Dashboard defines the directory structure on NAS where the flow records are exported.

You can export the flow records in Base or Full mode. In Base mode, only 5-tuple data for the flow record is exported. In Full mode the entire data for the flow record is exported.

Nexus Dashboard requires read and write permission to NAS in order to export the flow record. A system issue is raised if Nexus Dashboard fails to write to NAS.

Guidelines and limitations for network attached storage

- In order for Nexus Dashboard to export the flow records to an external storage, the Network Attached Storage added to Nexus Dashboard must be exclusive for Nexus Dashboard.
- Network Attached Storage with Network File System (NFS) version 3 must be added to Nexus Dashboard.
- Flow Telemetry and Netflow records can be exported.
- Export of FTE is not supported.
- Average Network Attached Storage requirements for 2 years of data storage at 20k flows per sec:
 - Base Mode: 500 TB data
 - Full Mode: 2.8 PB data
- If there is not enough disk space, new records will not be exported and an anomaly is generated.

Add network attached storage to export flow records

The workflow to add Network Attached Storage (NAS) to export flow records includes the following steps:

1. Add NAS to Nexus Dashboard.
2. Add the onboarded NAS to Nexus Dashboard to enable export of flow records.

Add NAS to Nexus Dashboard

Follow these steps to add NAS to Nexus Dashboard.

1. Navigate to **Admin > System Settings > General**.
2. In the **Remote storage** area, click **Edit**.
3. Click **Add Remote Storage Locations**.
4. Complete the following fields to add NAS to Nexus Dashboard.
 - a. Enter the name of the Network Attached Storage and a description, if desired.
 - b. In the **Remote storage location type** field, click **NAS Storage**.
 - c. In the **Type** field, choose **Read Write**.

Nexus Dashboard requires read and write permission to export the flow record to NAS. A system issue is raised if Nexus Dashboard fails to write to NAS.

- d. In the **Hostname** field, enter the IP address of the Network Attached Storage.
- e. In the **Port** field, enter the port number of the Network Attached Storage.

- f. In the **Export path** field, enter the export path.

Using the export path, Nexus Dashboard creates the directory structure in NAS for exporting the flow records.

- g. In the **Alert threshold** field, enter the alert threshold time.

Alert threshold is used to send an alert when the NAS is used beyond a certain limit.

- h. In the **Limit (Mi/Gi)** field, enter the storage limit in Mi/Gi.

- i. Click **Save**.



In a three-node Nexus Dashboard cluster deployment, NAS backup operations require that all cluster nodes are in an operational and healthy state. If any node is down, unreachable, or not in a fully functional state, NAS backup operations do not proceed.

This ensures cluster consistency and data integrity during backup operations. Verify cluster health and confirm that all three nodes are up and synchronized before initiating or scheduling NAS backups.

Add the onboarded NAS to Nexus Dashboard

Follow these steps to add the onboarded NAS to Nexus Dashboard.

1. Navigate to the Fabrics page:

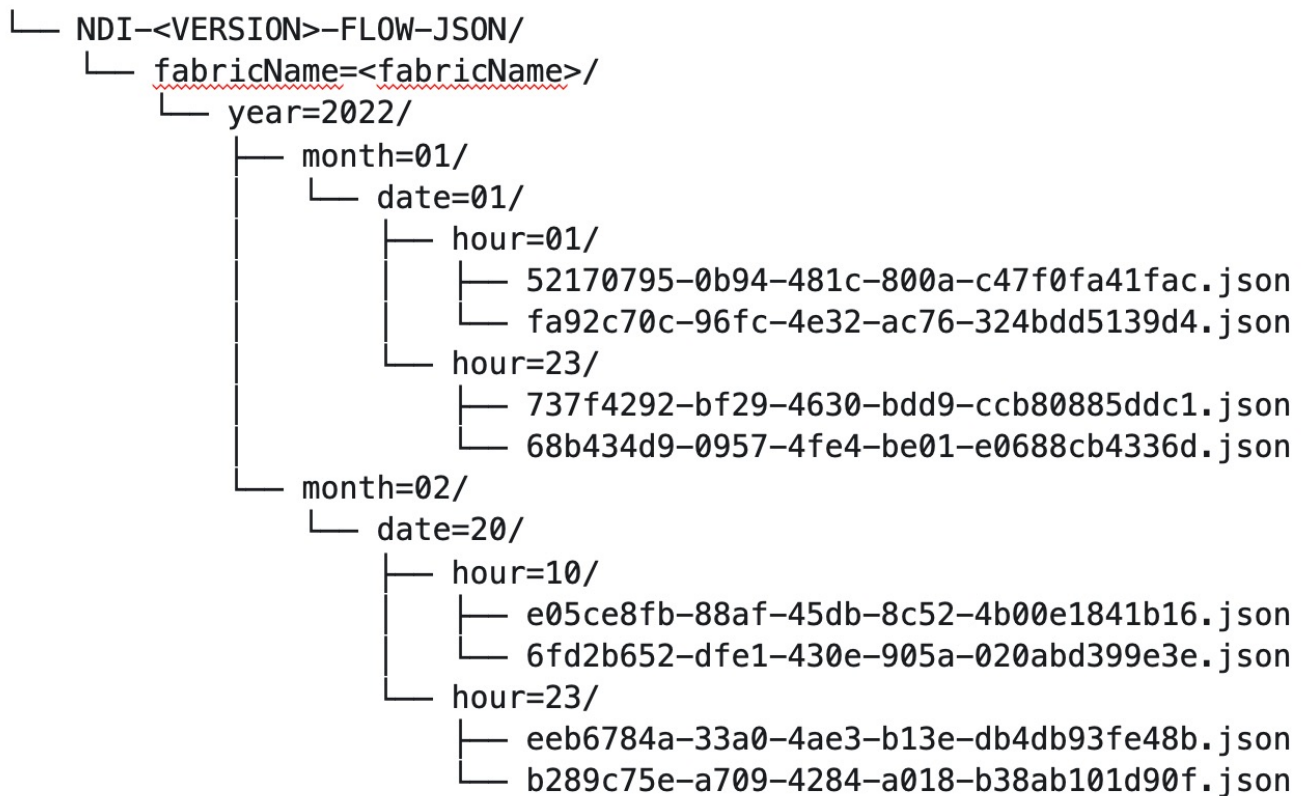
Manage > Fabrics

2. Choose the fabric with the telemetry feature enabled.
3. Choose **Actions > Edit Fabric Settings**.
4. Click **Telemetry**.
5. Click the **NAS** tab in the **Telemetry** page.
6. Make the necessary configurations in the **General settings** area.
 - a. Enter the name in the **Name** field.
 - b. In the **NAS server** field, choose the NAS server added to Nexus Dashboard from the drop-down list.
7. In the **Collection settings** area, choose the flow from the **Flows** drop-down list.
 - o In Base mode, only 5-tuple data for the flow record is exported.
 - o In Full mode, the entire data for the flow record is exported.
8. Click **Save**.

The traffic from the flows displayed in the **Flows** page is exported as a JSON file to the external NAS in the following directory hierarchy.



Beginning with Nexus Dashboard 4.3.1, flows are exported to a path with the unified Nexus Dashboard version in the directory name instead of the NIR version.



Navigate to **Analyze > Flows** to view the flows that will be exported.

Each flow record is written as a line delimited JSON.

JSON output file format for a flow record in base mode

```

{"fabricName":"myapic","terminalTs":1688537547433,"originTs":1688537530376,"srcIp":
:"2000:201:1:1::1","dstIp"::"2000:201:1:1::3","srcPort":1231,"dstPort":1232,"ingressVrf"
:"vrf1","egressVrf"::"vrf1","ingressTenant"::"FSV1","egressTenant"::"FSV1","protocol"::"U
DP"}

{"fabricName":"myapic","terminalTs":1688537547378,"originTs":1688537530377,"srcIp"
:"201.1.1.127","dstIp"::"201.1.1.1","srcPort":0,"dstPort":0,"ingressVrf"::"vrf1","egressVrf
"::"","ingressTenant"::"FSV2","egressTenant"::"","protocol"::"ANY-HOST"}
  
```

JSON output file format for a flow record in full mode

```

{"fabricName":"myapic","terminalTs":1688538023562,"originTs":1688538010527,"srcIp"
:"201.1.1.121","dstIp"::"201.1.1.127","srcPort":0,"dstPort":0,"ingressVrf"::"vrf1","egress
Vrf"::"vrf1","ingressTenant"::"FSV2","egressTenant"::"FSV2","protocol"::"ANY-
HOST","srcEpg"::"ext-epg","dstEpg"::"ext-
epg1","latencyMax":0,"ingressVif"::"eth1/15","ingressVni":0,"latency":0,"ingressNodes":
"Leaf1-
2","ingressVlan":0,"ingressByteCount":104681600,"ingressPktCount":817825,"ingressBur
st":0,"ingressBurstMax":34768,"egressNodes"::"Leaf1-2","egressVif"::"po4",
  
```

```
" egressVni":0," egressVlan":0," egressByteCount":104681600," egressPktCount":817825,"
 egressBurst":0," egressBurstMax":34768," dropPktCount":0," dropByteCount":0," dropCode
 ":" ", " dropScore":0," moveScore":0," latencyScore":0," burstScore":0," anomalyScore":0,"
 hashCollision":false," dropNodes":" [ ]", " nodeNames":" [\" Leaf1-
 2\"]", " nodeIngressVifs":" [\" Leaf1-2,eth1/15\"]", " nodeEgressVifs":" [\" Leaf1-2,po4\"]“
 , " srcMoveCount":0," dstMoveCount":0," moveCount":0," prexmit":0," rtoOutside":false," ev
 ents":" [[\\\" 1688538010527,Leaf1-2,0,3,1,no,no,eth1/15,,po4,po4,,,,,0,64,0,,,,,,\\\" ]]" }
```

Flow collection

Understanding flow telemetry

Flow telemetry allows users to see the path taken by different flows in detail. It also allows you to identify the EPG and VRF instance of the source and destination. You can see the switches in the flow with the help of flow table exports from the nodes. The flow path is generated by stitching together all the exports in order of the flow.

You can configure the Flow Telemetry rule for the following interface types:

- VRF instances
- Physical interfaces
- Port channel interfaces
- Routed sub-interfaces (Cisco ACI fabric)
- SVIs (Cisco ACI fabric)



In a Cisco ACI fabric, if you want to configure routed sub-interfaces from the UI, select L3 Out.

In an NX-OS fabric, physical or port channel flow rules are supported only on routed interfaces.

Flow telemetry monitors the flow for each fabric separately, as there is no stitching across the fabrics in a fabric group. Therefore, flow telemetry is for individual flows. For example, if there are two fabrics (fabric A and fabric B) within a fabric group, and traffic is flowing between the two fabrics, they will be displayed as two separate flows. One flow will originate from Fabric A and display where the flow exits. And the other flow from Fabric B will display where it enters and where it exits.

Flow telemetry guidelines and limitations

- All flows are monitored as a consolidated view in a unified pipeline for Cisco ACI and NX-OS fabrics, and the flows are aggregated under the same umbrella.
- Even if a particular node (for example, a third-party switch) is not supported for Flow Telemetry, Nexus Dashboard will use LLDP information from the previous and next nodes in the path to identify the switch name and the ingress and egress interfaces.
- You can enable Flow Telemetry at the fabric level in Nexus Dashboard, even if only a subset of switches in the fabric supports the feature. Enabling Flow Telemetry at the fabric level does not

require all switches to support it; provided at least one switch supports Flow Telemetry, you can enable the feature for the entire fabric.

- Nexus Dashboard supports Kafka export for Flow anomalies. However, Kafka export is not currently supported for Flow Event anomalies.

Flow telemetry guidelines and limitations for NX-OS fabrics

- Ensure that you have configured NTP and enabled PTP in Nexus Dashboard. See [Cisco Nexus Dashboard Deployment Guide](#) and [Precision Time Protocol \(PTP\) for Cisco Nexus Dashboard Insights](#) for more information. You are responsible for configuring the switches with external NTP servers.
- The **Fabric Overview** page displays the correct PTP status only when **Flow telemetry** or **Traffic analytics** is enabled under **Flow collection**. Enabling PTP in the fabric settings alone does not enable PTP monitoring.
- In the **Edit Flow** page, you can enable all three telemetry types. sFlow is most restrictive, Netflow has some more capability, and Flow Telemetry has the most capability. We recommend that you enable Flow Telemetry if it is available for your configuration. If Flow Telemetry is not available, then use Netflow. If Netflow is not available, use sFlow.
- If there are multiple Nexus Dashboard clusters onboarded to Nexus Dashboard, partial paths will be generated for each fabric.
- For inter-fabric flows that traverse an external fabric and return to the source fabric (a U-turn path), Nexus Dashboard displays the egress interface as **N/A** when the external fabric contains Cisco N9K-C9364C switches. Nexus Dashboard cannot capture egress port metadata for these flows. The remote port value remains valid.
- If you manually configure the fabric to use with Nexus Dashboard and Flow Telemetry support, the Flows Exporter port changes from 30000 to 5640. To prevent a breakage of the Flows Exporter, adjust the automation.
- Flow telemetry is supported in -FX3 platform switches for the following NX-OS versions:
 - 9.3(7) and later
 - 10.1(2) and later
 - Flow telemetry is not supported in -FX3 platform switches for NX-OS version 10.1(1).
- Interface based Flow Telemetry is only supported on modular chassis with -FX land -GX line cards on physical ports and port-channels rules.
- If interface-based Flow Telemetry is pushed from Nexus Dashboard for **Classic LAN** and **External Connectivity Network** fabrics, perform the following steps:
 - Choose the fabric.
 - Choose **Policies** > **Action** > **Add policy** > **Select all** > **Choose template** > host_port_resync and click **Save**.
 - In the Fabric Overview page, choose **Actions** > **Recalculate and deploy**.
- For VXLAN fabrics, interface-based Flow Telemetry is not supported on switch links between spine switch and leaf switch.
- If you want to use the default VRF instance for flow telemetry, you must create the VRF instance with a name of "default" in lowercase. Do not enter the name with any capital letters.

- Flow telemetry is not supported in classic LAN topologies with 2-level VPC access layers.
- If you want to enable Flow Telemetry, ensure that there are no pre-existing Netflow configurations on the switches. If there are any pre-existing configurations, the switch configuration may fail.

To enable Flow Telemetry without configuration issues, follow these steps:

- Ensure that there are no pre-existing Netflow configurations on the switches. If such configurations exist, enabling Flow Telemetry might result in a system anomaly with an error message stating **invalid command match IP source address**.
- If you encounter the error, disable Flow Telemetry.
- Remove any existing Netflow configurations from the switches.
- Re-enable Flow Telemetry.
- For some flows, latency information is not available, which could happen due to latency issues. In these cases, latency information will be reported as 0.

Flow telemetry rules guidelines and limitations for NX-OS fabrics

- If you configure an interface rule (physical or port channel) on a subnet, it can monitor only incoming traffic. It cannot monitor outgoing traffic on the configured interface rule.
- If a configured port channel that contains two physical ports, only the port channel rule is applicable. Even if you configure physical interface rules on the port, only port channel rule takes precedence.
- For NX-OS release 10.3(2) and earlier, if a flow rule are configured on an interface, then global flow rules are not matched.
- For NX-OS release 10.3(3) and later, a flow rule configured on an interface is matched first and then the global flow rules are matched.

Configure flows

Configure flow collection modes

Follow these steps to configure flow collection modes.

1. Navigate to **Admin > System Settings > Flow collection**.
2. In the **Flow collection mode** area, choose **Flow telemetry**.



Enabling Flow Telemetry automatically activates Flow Telemetry Events. Whenever a compatible event takes place, an anomaly will be generated, and the What's the impact? section in the **Anomaly** page will display the associated flows. You can manually configure a Flow Telemetry rule to acquire comprehensive end-to-end information about the troublesome flow.

Configure flow collection rules in an NX-OS fabric

Follow these steps to configure flow collection rules in an NX-OS fabric.

1. Navigate to the **Telemetry** page for your fabric.
 - a. Navigate to the main **Fabrics** page:

Manage > Fabrics

- b. 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 telemetry settings.
- c. Single-click on that fabric.

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

- d. Click **Actions > Edit Fabric Settings**.

The **Edit *fabric_name* Settings** page appears.

- e. Verify that the **Telemetry** option is enabled in the **Enabled features** area.

The Telemetry tab doesn't become available unless the **Telemetry** option is enabled in the **Enabled features** area.

- f. Click the **Telemetry** tab to access the telemetry settings for this fabric.

2. Click the **Flow collection** tab in the **Telemetry** page.
3. In the **Mode** area, click **Flow telemetry**.
4. In the **Flow collections rules** area, determine what sort of flow collection rule that you want to add.
 - o [VRF](#)
 - o [Physical interface](#)
 - o [Port channel](#)

VRF

Follow these steps to add a VRF rule.

1. Click the **VRF** tab.

A table with already-configured VRF flow collection rules is displayed.

For any VRF flow collection rule in this table, click the ellipsis (...), then click **Edit rule** to edit that rule or **Delete rule** to delete it.

2. Add a new rule by clicking **Create flow collection rule**.
 - a. In the **General** area, complete the following:
 - i. Enter the name of the rule in the **Rule Name** field.
 - ii. The VRF field is disabled. The flow rule applies to all the VRF instances.
 - iii. In the **Flow Properties** area, select the protocol for which you intend to monitor the flow traffic.
 - iv. Enter the source and destination IP addresses. Enter the source and destination port.
 - v. Click **Save**.

Physical interface

Follow these steps to add a physical interface rule.

1. Click the **Physical interface** tab.

A table with already-configured physical interface flow collection rules is displayed.

For any physical interface flow collection rule in this table, click the ellipsis (...), then click **Edit rule** to edit that rule or **Delete rule** to delete it.

2. Add a new rule by clicking **Create flow collection rule**.

- a. In the **General** area, complete the following:

- i. Enter the name of the rule in the **Rule Name** field.
- ii. Check the **Enabled** check box to enable the status. If you enable the status, the rule will take effect. Otherwise, the rule will be removed from the switches.
- iii. In the **Flow Properties** area, select the protocol for which you intend to monitor the flow traffic.
- iv. Enter the source and destination IP addresses. Enter the source and destination port.
- v. In the **Interface List** area, click **Select a Node**. Use the search box to select a node.
- vi. From the drop-down list, select an interface. You can add more than one row (node+interface combination) by clicking **Add Interfaces**. However, within the rule, a node can appear only once. Configuration is rejected if more than one node is added.
- vii. Click **Save**.

Port channel

Follow these steps to add a port channel rule.

1. Click the **Port channel** tab.

A table with already-configured port channel flow collection rules is displayed.

For any port channel flow collection rule in this table, click the ellipsis (...), then click **Edit rule** to edit that rule or **Delete rule** to delete it.

2. Add a new rule by clicking **Create flow collection rule**.

- a. In the **General** area, enter the name of the rule in the **Rule Name** field.

- i. Select the **Enabled** check box to enable the status. If you enable the status, the rule will take effect. Otherwise, the rule will be removed from the switches.
- ii. In the **Flow Properties** area, select the protocol for which you intend to monitor the flow traffic.
- iii. Enter the source and destination IP addresses. Enter the source and destination port.
- iv. From the drop-down list, select an interface. You can add more than one row (node+interface combination) by clicking **Add Interfaces**. However, within the rule, a node can appear only once. Configuration is rejected if more than one node is added.
- v. Click **Save**.

3. Click **Done**.

Monitor the subnet for flow telemetry

In the following example, the configured rule for a flow monitors the specific subnet provided. The rule is pushed to the fabric which pushes it to the switches. So, when the switch sees traffic coming from a source IP or the destination IP, and if it matches the subnet, the information is captured in the TCAM and exported to the Nexus Dashboard service. If there are 4 nodes (A, B, C, D), and the traffic moves from A > B > C > D, the rules are enabled on all 4 nodes and the information is captured by all the 4 nodes. Nexus Dashboard stitches the flows together. Data such as the number of drops and the number of packets, anomalies in the flow, and the flow path are aggregated for the 4 nodes.

Follow these steps to monitor the subnet for flow telemetry.

1. Navigate to **Manage > Fabric**.
2. Choose a fabric.
3. Verify that your **Fabrics** and the **Snapshot** values are appropriate. The default snapshot value is 15 minutes. Your choice will monitor all the flows in the chosen fabric or snapshot fabric.
4. Navigate to **Connectivity > Flows** to view a summary of all the flows that are being captured based on the snapshot that you chose.

The related anomaly score, record time, the nodes sending the flow telemetry, flow type, ingress and egress nodes, and additional details are displayed in a table format. If you click a specific flow in the table, specific details are displayed in the sidebar for the particular flow telemetry. In the sidebar, if you click the Details icon, the details are displayed in a larger page. In this page, in addition to other details, the **Path Summary** is also displayed with specifics related to source and destination. If there are flows in the reverse direction, that will also be visible in this location.

For a bi-directional flow, there is an option to choose to reverse the flow and see the path summary displayed. If there are any packet drops that generate a flow event, they can be viewed in the Anomaly dashboard.

Understanding Netflow

Netflow is an industry standard where Cisco routers monitor and collect network traffic on an interface. Netflow version 9 is supported.

Netflow enables the network administrator to determine information such as source, destination, class of service, and causes of congestion. Netflow is configured on the interface to monitor every packet on the interface and provide telemetry data. You cannot filter on Netflow.

Netflow in Nexus series switches is based on intercepting the packet processing pipeline to capture summary information of network traffic.

The components of a flow monitoring setup are as follows:

- **Exporter:** Aggregates packets into flows and exports flow records towards one or more collectors
- **Collector:** Reception, storage, and pre-processing of flow data received from a flow exporter
- **Analysis:** Used for traffic profiling or network intrusion

- The following interfaces are supported for Netflow:

Supported interfaces for Netflow

Interfaces	5 Tuple	Nodes	Ingress	Egress	Path	Comments
Routed Interface/Port Channel	Yes	Yes	Yes	No	Yes	Ingress node is shown in path
Sub Interface/Logical (Switch Virtual Interface)	Yes	Yes	No	No	No	No



In an NX-OS fabric, port channel support is available if you monitor only the host-facing interfaces.

Understanding Netflow types

You can use these Netflow types.

- [Full Netflow](#)
- [Sampled Netflow](#)

Full Netflow

With Full Netflow, all packets on the configured interfaces are captured into flow records in a flow table. Flows are sent to the supervisor module. Records are aggregated over configurable intervals and exported to the collector. Except in the case of aliasing (multiple flows hashing to the same entry in the flow table), all flows can be monitored regardless of their packet rate.

Cisco N9000 Series switches with the Fabric Controller type as well as switches in a Cisco ACI fabric support Full Netflow.

Sampled Netflow

With Sampled Netflow, packets on configured interfaces are time sampled. Flows are sent to the supervisor or a network processor for aggregation. Aggregated flow records are exported at configured intervals. The probability of a record for a flow being captured depends on the sampling frequency and packet rate of the flow relative to other flows on the same interface.

Nexus 7000 and Nexus 7700 Series switches with F/M line cards and the Fabric Controller type, support Sampled Netflow.

Netflow guidelines and limitations

- In Cisco N9000 series switches, Netflow supports a small subset of the published export fields in the RFC.
- Netflow is captured only on the ingress port of a flow as only the ingress switch exports the flow. Netflow cannot be captured on fabric ports.

Netflow guidelines and limitations for Cisco ACI fabrics

- We recommend that you enable Flow Telemetry. If that is not available for your configuration, use Netflow. However, you can determine which mode of flow to use based upon your fabric configuration.
- Enabling both Flow Telemetry and Netflow is not supported.
- After you enable Netflow, you must obtain the Netflow collector IP address and configure Cisco APIC with the collector IP address. See [Cisco APIC and NetFlow](#).

To obtain the Netflow collector IP address, navigate to **Admin > System Settings > General**, then locate the **External Pools** area. Click **View all** at the bottom left area of the **External Pools** tile; the **Telemetry-collector** persistent IP addresses listed in the table are used for the Netflow collector IP address.

- The Netflow and sFlow flow collection modes do not support any anomaly.

Netflow guidelines and limitations for NX-OS fabrics

- In the **Edit Flow** page, you can enable all three modes. Choose the best possible mode for a product. sFlow is the most restrictive, Netflow has more capabilities, and Flow Telemetry has the most capabilities. We recommend that you enable Flow Telemetry if it is available for your configuration. If Flow Telemetry is not available, then use Netflow. If Netflow is not available, use sFlow.
- In Cisco Nexus 7000 and Cisco N9000 Series switches, only the ingress host-facing interface configured for Netflow are supported (either in VXLAN or Classic LAN).
- The Netflow supported fabrics are Classic and VXLAN. VXLAN is not supported on fabric ports.
- Netflow configurations will not be pushed. However, if a fabric is managed, the software sensors will be pushed.
- If you manually configure the fabric to use with Nexus Dashboard and Netflow support, the Flows Exporter port changes from 30000 to 5640. To prevent a breakage of the Flows Exporter, adjust the automation.
- To configure Netflow on fabric switches, see the **Configuring Netflow** section in the [Cisco N9000 Series NX-OS System Management Configuration Guide](#).

Configure Netflow

Follow these steps to configure Netflow.

1. Navigate to the **Telemetry** page for your LAN or IPFM fabric.
2. Click the **Flow collection** tab on the **Telemetry** page.
3. In the **Mode** area, make the following choices:
 - Choose **Netflow**.
 - Choose **Flow Telemetry**.
4. Click **Save**.

Understanding sFlow

sFlow is an industry standard technology traffic in data networks containing switches and routers. Nexus Dashboard supports [sFlow version 5](#) on Cisco Nexus 3000 series switches.

sFlow provides the visibility to enable performance optimization, an accounting and billing for usage, and defense against security threats.

The following interfaces are supported for sFlow:

Supported interfaces for sFlow

Interfaces	5 Tuple	Nodes	Ingress	Egress	Path	Comments
Routed Interface	Yes	Yes	Yes	Yes	Yes	Ingress node is shown in path

Guidelines and limitations for sFlow

- Nexus Dashboard supports sFlow with Cisco Nexus 3000 series switches.
- It is recommended to enable Flow Telemetry if it is available for your configuration. If it is not available for your configuration, use Netflow. If Netflow, is not available for your configuration, then use sFlow.
- For sFlow, Nexus Dashboard requires the configuration of persistent IPs under cluster configuration, and 6 IPs in the same subnet as the data network are required.
- sFlow configurations will not be pushed. However, if a fabric is managed, the software sensors will be pushed.
- If you manually configure the fabric to use with Nexus Dashboard and sFlow support, the Flows Exporter port changes from 30000 to 5640. To prevent a breakage of the Flows Exporter, adjust the automation.
- Nexus Dashboard does not support sFlow in the following Cisco Nexus 3000 Series switches:
 - Cisco Nexus 3600-R Platform Switch (N3K-C3636C-R)
 - Cisco Nexus 3600-R Platform Switch (N3K-C36180YC-R)
 - Cisco Nexus 3100 Platform Switch (N3K-C3132C-Z)
- Nexus Dashboard does not support sFlow in the following Cisco N9000 Series fabric modules:
 - Cisco N9508-R fabric module (N9K-C9508-FM-R)
 - Cisco N9504-R fabric module (N9K-C9504-FM-R)
- To configure sFlow on fabric switches, see the **Configuring sFlow** section in the [Cisco N9000 Series NX-OS System Management Configuration Guide](#).

Configure sFlow telemetry

Prerequisites

Follow these steps to configure sFlow telemetry.

1. Navigate to the **Telemetry** page for your LAN or IPFM fabric.
2. Click the **Flow collection** tab on the **Telemetry** page.
3. In the **Mode** area, make the following choices:
 - Choose **sFlow**.
 - Choose **Flow Telemetry**.
4. Click **Save**.

External streaming

The **External streaming** tab in Nexus Dashboard allows you export data that Nexus Dashboard collects over Kafka, email, and syslog. Nexus Dashboard generates data such as advisories, anomalies, audit logs, faults, statistical data, and risk and conformance reports. When you configure a Kafka broker, Nexus Dashboard writes all data to a topic. By default, the Nexus Dashboard collects export data every 30 seconds or at a less frequent interval.

For ACI fabrics, you can also collect data for specific resources (CPU, memory, and interface utilization) every 10 seconds from the leaf and spine switches using a separate data pipeline. To export this data, select the **Usage** option under **Collection Type** in the **Message bus** export settings. Additionally, CPU and memory data is collected for the controllers.



Nexus Dashboard does not store the collected data in Elasticsearch; instead, it exports the data directly to your repository or data lake using a Kafka broker for consumption. By using the Kafka export functionality, you can then export this data to your Kafka broker and push it into your data lake for further use.

You can configure an email scheduler to define the type of data and the frequency at which you want to receive information via email. You can also export anomaly records to an external syslog server. To do this, select the **Syslog** option under the **External Streaming** tab.

Configure external streaming settings

Follow these steps to configure external streaming settings.

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

Go to **Manage > Fabrics**.

2. Choose the fabric for which you configure streaming settings.
3. From the **Actions** drop-down list, choose **Edit fabric settings**.

The **Edit *fabric-name* settings** page displays.



You can also access the **Edit *fabric-name* settings** page for a fabric from the **Fabric Overview** page. In the **Fabric Overview** page, click the **Actions** drop-down list and choose **Edit fabric settings**.

4. In the **Edit *fabric-name* settings** page, click the **External streaming** tab.

You can view these options.

- o Email
- o Message bus
- o Syslog

Guidelines and limitations

- Intersight connectivity is required to receive the reports by email.
- You can configure up to five emails per day for periodic job configurations.
- A maximum of six exporters is supported for export across all types of exporters including email, message bus, and syslog. You must provide unique names for each export.
- The scale for Kafka exports is increased to support up to 20 exporters per cluster. However, statistics selection is limited to any six exporters.
- Before configuring your Kafka export, you must add the external Kafka IP address as a known route in your Nexus Dashboard cluster configuration and verify that Nexus Dashboard can reach the external Kafka IP address over the network.
- The anomalies in Kafka and email messages include categories such as Resources, Environmental, Statistics, Endpoints, Flows, and Bugs.
- Export data is not supported for snapshot fabrics.
- You must provide unique names for each exporter, and they may not be repeated between Kafka export for **Alerts and Events** and Kafka export for **Usage**.
- Nexus Dashboard supports Kafka export for flow anomalies. However, Kafka export is not currently supported for flow Event anomalies.

Guidelines and limitations in NX-OS fabrics

- Remove all configurations in the *Message Bus Configuration* and *Email* page before you disable Software Telemetry on any fabric and remove the fabric from Nexus Dashboard.
- If you use interface groups while enabling PTP in an Enhanced Classic LAN fabric, you must add the following commands to the freeform section of the interface group.
 - **ttag**
 - **ttag-strip**

If you do not add these commands, endpoints connected behind these interfaces will experience traffic loss.

Email

The email scheduler feature in Nexus Dashboard automates the distribution of summarized data collected from Nexus Dashboard. It allows customization of selection of email recipients, choice of email format, scheduling frequency settings, and configuring the types of alerts and reports.

Follow these steps to configure an email scheduler.

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

Go to **Manage > Fabrics**.

2. Choose the fabric for which you configure streaming settings.
3. From the **Actions** drop-down list, choose **Edit fabric settings**.

The **Edit *fabric-name* settings** page displays.

4. In the **Edit *fabric-name* settings** page, click the **External streaming** tab.
5. Click the **Email** tab.
6. Review the information provided in the **Email** tab for already-configured email configurations.

The following details display under **Email** tab.

Field	Description
Name	The name of the email configuration.
Email	The email addresses used in the email configuration.
Start time	The start date used in the email configuration.
Frequency	The frequency in days or weeks set in the email configuration.
Anomalies	The severity level for anomalies and advisories set in the email configuration.
Advisories	
Risk and conformance reports	The status of the overall inventory for a fabric, including software release, hardware platform, and a combination of software and hardware conformance.

To add a new email configuration, click **Add email** in the **Email** page.

1. Follow these steps to configure **General Settings**.
 - a. In the **Name** field, enter the name of the email scheduler.
 - b. In the **Email** field, enter one or more email addresses separated by commas.
 - c. In the **Format** field, choose **Text** or **HTML** email format.
 - d. In the **Start date** field, choose the start date when the scheduler should begin sending emails.
 - e. In the **Collection frequency in days** field, specify how often the summary is sent, you can choose days or weeks.

General settings

Name *

Email *

Use commas to enter multiple email addresses

Format

☒ Text ☐ HTML

Start date *

Collection frequency in days *

 Days 

2. Follow these steps to configure **Collection Settings**.

- a. In the **Mode** field, choose one of the following modes.

- **Basic**—displays the severity levels for anomalies and advisories.
 - **Advanced**—displays the categories and severity levels for anomalies and advisories.
- b. Check the **Only include active alerts in email** check box, to include only active anomaly alerts.
 - c. Under **Anomalies** choose the categories and severity levels for the anomalies.
 - d. Under **Advisories** choose the categories and severity levels for the advisories.
 - e. Under **Risk and Conformance Reports**, choose from the following options.
 - Software
 - Hardware

Collection settings

Mode
☒ Basic ☐ Advanced

Only include active alerts in email
☐

Anomalies [Select all](#) [Clear all](#)

☐ Critical
☐ Major
☐ Warning
☐ Minor

Advisories [Select all](#) [Clear all](#)

☐ Critical
☐ Major
☐ Warning
☐ Minor

Risk and Conformance Reports [Select all](#) [Clear all](#)

☐ Software
☐ Hardware

[Cancel](#) [Save](#)

3. Click **Save**.

The **Email** area displays the configured email schedulers.

You will receive an email about the scheduled job on the provided *Start Date* and at the time provided in the **Collection frequency in days** field. The subsequent emails follow after *Collect Every* frequency expires. If the provided time is in the past, Nexus Dashboard will send you an email immediately and trigger the next email after the duration from the provided start time expires.

Message bus

Add Kafka broker configuration

Follow these steps to configure the message bus and add kafka broker.

1. Configure the message bus at the **System Settings** level.
 - a. Navigate to **Admin > System Settings > General**.
 - b. In the **Message bus configuration** area, click **Edit**.

The **Message bus configuration** dialog box opens.

- c. Click **Add message bus configuration**.

The **Add message bus configuration** dialog box opens.

- d. In the **Name** field, enter a name for the configuration.
- e. In the **Hostname/IP address** and **Port** fields, enter the IP address of the message bus consumer and the port that is listening on the message bus consumer.
- f. In the **Topic name** field, enter the name of the Kafka topic to which Nexus Dashboard must send the messages.
- g. In the **Mode** field, choose the security mode.

The supported modes are **Unsecured**, **Secured SSL** and **SASLPLAIN**. The default value is **Unsecured**.

- For **Unsecured**, no other configurations are needed.
- For **Secured SSL**, fill out the following field:

Client certification name—The System Certificate name configured at the Certificate Management level. The CA certificate and System Certificate (which includes Client certificate and Client key) are added at the Certificate Management level.

Refer to Step 2 for step-by-step instructions on managing certificates. Navigate to **Admin > Certificate Management** to manage the following certificates:

- **CA Certificate**—The CA certificate used for signing consumer certificate, which will be stored in the trust-store so that Nexus Dashboard can trust the consumer.
 - **Client Certificate**—The CA signed certificate for Nexus Dashboard. The certificate is signed by the same CA, and the same CA certificate will be in the truststore of the consumer. This will be stored in Nexus Dashboard's Kafka keystore that is used for exporting.
 - **Client Key**—A private key for the Kafka producer, which is Nexus Dashboard in this case. This will be stored in Nexus Dashboard's Kafka keystore that is used for exporting.
- For **SASLPLAIN**, fill out these fields:
 - **Username**—The username for the SASL/PLAIN authentication.
 - **Password**—The password for the SASL/PLAIN authentication.

- h. Click **Save**

2. Add CA certificates and System certificates at the **Certificate Management level**.

- a. Navigate to **Admin > Certificate Management**.
- b. In the **Certificate management** page, click the **CA Certificates** tab, then click **Add CA certificate**.

The fields in the **CA Certificates** tab are described in the following table.

Field	Description
Certificate name	The name of the CA certificate.
Certificate details	The details of the CA certificate.
Attached to	The CA signed certificate attached to Nexus Dashboard.
Expires on	The Expiry date and time of the CA certificate.
Last updated time	The last updated time of the CA certificate.

- c. In the **Certificate management** page, click the **System certificates** tab, then click **Add system certificate** to add Client Certificate and Client key. Note that the Client certificate and Client key should have same names except extensions as .cer/.crt/.pem for Client certificate and .key for Client key.



You must add a valid CA Certificate before adding the corresponding System Certificate.

The fields in the **System Certificates** tab are described in the following table.

Field	Description
Certificate name	The name of the Client certificate.
Certificate details	The details of the Client certificate.
Attached to	The feature to which the system certificate is attached to, in this case, the message bus.
Expires on	The Expiry date and time of the CA certificate.
Last updated time	The last updated time of the CA certificate.



To configure message bus, the System Certificate should be attached to message bus feature.

To attach a System Certificate to the message bus feature:

- Choose the System Certificate that you want to use and click the ellipses (...) on that row.
- Choose **Manage Feature Attachments** from the drop-down list.

The **Manage Feature Attachments** dialog box opens.

- In the **Features** field, choose **messageBus**.
- Click **Save**.

For more information on CA certificates, see [Managing Certificates in your Nexus Dashboard](#).

Configure Kafka exports in fabric settings

- Navigate to the **External streaming** page for your fabric.
 - Navigate to the **Fabrics** page.

Go to **Manage > Fabrics**.

- b. Choose the fabric for which you configure streaming settings.
- c. From the **Actions** drop-down list, choose **Edit fabric settings**.

The **Edit fabric-name settings** page displays.

- d. In the **Edit fabric-name settings** page, click the **External streaming** tab.
 - e. Click the **Message bus** tab.
2. Review the information provided in the **Message bus** tab for already-configured message bus configurations, or click **Add message bus** to add a new message bus configuration.

Skip to Step 3 if you are adding a message bus.

The fields in the **Message bus** tab are described in the following table.

Field	Description
Message bus stream	The name of the message bus stream configuration.
Collection type	The collection type used by the message bus stream.
Mode	The mode used by the message bus stream.
Anomalies	The severity level for anomalies and advisories set in the message bus stream configuration.
Advisories	
Statistics	The statistics that were configured for the message bus stream.
Faults	The severity level for faults set in the message bus stream configuration.
Audit Logs	The audit logs that were configured for the message bus stream.

3. To configure a new message bus stream, in the **Message bus** page, click **Add message bus**.
4. In the **Message bus stream** field, choose the message bus stream that you want to edit.
5. In the **Collection Type** area, choose the appropriate collection type.

Depending on the **Collection Type** that you choose, the options displayed in this area will change.

- **Alerts and events:** This is the default setting. Continue to Step 7, if you choose **Alerts and events**.
- **Usage:** In the **Collection settings** area, under **Data**, the Resources, and Statistics for the collection settings are displayed. By default, the data for CPU, Memory, and Interface Utilization are collected and exported. You cannot choose to export a subset of these resources.



Usage is applicable only for ACI Fabrics. This option is disabled for other fabrics.

6. Click **Save**. The configured message bus streams are displayed in the **Message bus** area. This configuration now sends immediate notification when the selected anomalies or advisories occur.
7. If you choose **Alerts and events** as the **Collection Type**, in the **Mode** area, choose either **Basic** or **Advanced**.

The configurations that are available in each collection settings section might vary, depending on

the mode that you set.

8. Determine which area you want to configure for the message bus stream.

The following areas appear in the page:

- [Anomalies](#)
- [Advisories](#)
- [Statistics](#)
- [Faults](#)
- [Audit Logs](#)

After you complete the configurations on this page, click **Save**. Nexus Dashboard displays the configured message bus streams in the **Message bus** area. This configuration now sends immediate notification when the selected anomalies or advisories occur.

Anomalies

- If you chose **Basic** in the **Mode** area, choose one or more of the following severity levels for anomaly statistics that you want to configure for the message bus stream:

- Critical
- Major
- Warning
- Minor

Or click **Select all** to select all available statistics for the message bus stream.

- If you chose **Advanced** in the **Mode** area:
 - Choose one or more of the following categories for anomaly statistics that you want to configure for the message bus stream:
 - Active Bugs
 - Capacity
 - Compliance
 - Configuration
 - Connectivity
 - Hardware
 - Integrations
 - System
 - Choose one or more of the following severity levels for anomaly statistics that you want to configure for the message bus stream:
 - Critical
 - Major
 - Warning

- Minor

Or click **Select all** to select all available categories and statistics for the message bus stream.
For more information on anomaly levels, see [Detecting Anomalies in Your Nexus Dashboard](#).

Advisories

- If you chose **Basic** in the **Mode** area, choose one or more of the following severity levels for advisory statistics that you want to configure for the message bus stream:

- Critical
- Major
- Warning
- Minor

Or click **Select all** to select all available statistics for the message bus stream.

- If you chose **Advanced** in the **Mode** area:
 - Choose one or more of the following categories for advisory statistics that you want to configure for the message bus stream:
 - Best Practices
 - Field Notices
 - HW end-of-life
 - SW end-of-life
 - PSIRT
 - Choose one or more of the following severity levels for advisory statistics that you want to configure for the message bus stream:
 - Critical
 - Major
 - Warning
 - Minor

Or click **Select all** to select all available categories and statistics for the message bus stream.
For more information on advisory levels, see [Identifying Advisories in Your Nexus Dashboard](#).

Statistics

There are no differences in the settings in the **Statistics** area when you choose **Basic** or **Advanced** in the **Mode** area.

Choose one or more of the following categories to configure the statistics you want to stream over the message bus.

- Interfaces—streams statistics related to network interfaces, such as traffic volume, error rates, and interface status.
- Protocol—streams protocol-specific statistics, including packet counts, protocol errors, and handshake success rates.

- Resource Allocation—streams data about system resources, such as CPU usage, memory consumption, and bandwidth allocation.
- Environmental—streams environmental metrics such as temperature, humidity, and power supply status.
- Endpoints—streams statistics about connected endpoints, including connection status, data throughput, and session durations.

Faults

There are no differences in the settings in the **Faults** area when you choose **Basic** or **Advanced** in the **Mode** area.

Choose one or more of the following severity levels for fault statistics that you want to configure for the message bus stream:

- Critical
- Major
- Minor
- Warning
- Info

Audit Logs

There are no differences in the settings in the **Audit Logs** area when you choose **Basic** or **Advanced** in the **Mode** area.

Choose one or more of the following categories for audit logs that you want to configure for the message bus stream:

- Creation
- Deletion
- Modification

Syslog

Nexus Dashboard supports the export of anomalies in syslog format. You can use the syslog configuration feature to develop network monitoring and analytics applications on top of Nexus Dashboard, integrate with the syslog server to get alerts, and build customized dashboards and visualizations.

After you choose the fabric where you want to configure the syslog exporter and set up the syslog export configuration, Nexus Dashboard establishes a connection with the syslog server and sends data to the syslog server.

Nexus Dashboard exports anomaly records to the syslog server. With syslog support, you can export anomalies to your third-party tools even if you do not use Kafka.

Guidelines and limitations for syslog

If the syslog server is not operational at a certain time, messages generated during that downtime will not be received by a server after the server becomes operational.

Add syslog server configuration

Follow these steps to add syslog server configuration.

1. Navigate to **Admin > System Settings > General**.
2. In the **Remote streaming servers** area, click **Edit**.

The **Remote streaming servers** page displays.

3. Click **Add server**.

The **Add server** page displays.

4. Choose the **Service** as **Syslog**.
5. Choose the **Protocol**.

You have these options.

- TCP
- UDP

6. In the **Name** field, provide the name for the syslog server.
7. In the **Hostname/IP address** field, provide the hostname or IP address of the syslog server.
8. In the **Port** field, specify the port number used by the syslog server.
9. If you want to enable secure communication, check the **TLS** check box.



Before you enable **TLS** you must upload the CA certificate for the syslog destination host to Nexus Dashboard. For more information see, [Upload a CA certificate](#).

Configure syslog to enable exporting anomalies data to a syslog server

Follow these steps to configure syslog to enable exporting anomalies data to a syslog server.

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

Go to **Manage > Fabrics**.

2. Choose the fabric for which you configure streaming settings.
3. From the **Actions** drop-down list, choose **Edit fabric settings**.

The **Edit *fabric-name* settings** page displays.

4. In the **Edit *fabric-name* settings** page, click the **External streaming** tab.
5. Click the **Syslog** tab.

The following details display under **Syslog** tab.

Edit CS1 Settings

Some fabric settings can only be modified from the ND Cluster owner standaloneM5PND246

General Telemetry **External streaming**

Email Message bus **Syslog**

General settings

Syslog server*

Select servers

Facility

Collection settings

Anomalies [Select all](#) [Clear all](#)

☐ Critical

☐ Major

☐ Warning

Cancel Save

6. Make the necessary configurations in the **General settings** area.

a. In the **Syslog server** drop down list, choose a syslog server.

The **Syslog server** drop down list displays the syslog servers that you added in the **System Settings** level. For more information, see [Add syslog server configuration](#).

b. In the **Facility** field, from the drop-down list, choose the appropriate facility string.

A facility code is used to specify the type of system that is logging the message. For this feature, the **local0-local7** keywords for locally used facility are supported.

7. In the **Collection settings** area, choose the desired severity options.

The options available are **Critical**, **Major**, and **Warning**.

8. Click **Save**.

Upload a CA certificate

Follow these steps to upload a CA certificate for syslog server **TLS**.

1. Navigate to **Admin > Certificate Management**.
2. In the **Certificate management** page, click the **CA certificates** tab, then click **Add CA certificate**.

You can upload multiple files at a single instance.

3. Browse your local directory and choose the certificate-key pair to upload.

You can upload certificates with the **.pem/.cer/.crt/** file extensions.

4. Click **Save** to upload the selected files to Nexus Dashboard.

A successful upload message appears. The uploaded certificates are listed in the table.

Additional settings

The following sections provide information for additional settings that might be necessary when editing the settings for an AI Data Center VXLAN fabric.

VXLAN EVPN fabrics provisioning

Cisco Nexus Dashboard provides an enhanced fabric workflow for unified underlay and overlay provisioning of the VXLAN BGP EVPN configuration on Cisco N9000 and Cisco Nexus 3000 series of switches. The configuration of the fabric is achieved via a powerful, flexible, and customizable template-based framework. Using minimal user inputs, an entire fabric can be brought up with Cisco-recommended best practice configurations in a short period of time. The set of parameters exposed in the Fabric Settings allow you to tailor the fabric to your preferred underlay provisioning options.

Border devices in a fabric typically provide external connectivity via peering with appropriate edge/core/WAN routers. These edge/core routers may either be managed or monitored by Nexus Dashboard. These devices are placed in a special fabric called the External Fabric. The same Nexus Dashboard can manage multiple VXLAN BGP EVPN fabrics while also offering easy provisioning and management of Layer-2 and Layer-3 DCI underlay and overlay configuration among these fabrics using a special construct called a VXLAN fabric group.

The Nexus Dashboard GUI functions for creating and deploying VXLAN BGP EVPN fabrics are as follows:

Manage > Fabrics > Create Fabric or Manage > Fabrics > Edit Fabric Settings

Create, edit, and delete a fabric:

- Create new VXLAN, VXLAN fabric group, and external VXLAN fabrics.
- View the VXLAN and VXLAN fabric group fabric topologies, including connections between fabrics.
- Update fabric settings.
- Save and deploy updated changes.
- Delete a fabric (if devices are removed).

Manage > Inventory

Device discovery and provisioning start-up configurations on new switches:

- Add switch instances to the fabric.
- Provision start-up configurations and an IP address to a new switch through POAP configuration.
- Update switch policies, save, and deploy updated changes.
- Create intra-fabric and inter-fabric links (also called Inter-Fabric Connections [IFCs]).

Manage > Fabrics > Connectivity > Interfaces > Actions > Create New Interface

Underlay provisioning:

- Create, deploy, view, edit, and delete a port-channel, vPC switch pair, Straight Through FEX (ST-

FEX), Active-Active FEX (AA-FEX), loopback, subinterface, etc.

- Create breakout and unbreakout ports.
- Shut down and bring up interfaces.
- Rediscover ports and view interface configuration history.

Manage > Inventory > Switches > Actions > Add Switches

Overlay network provisioning.

- Create new overlay networks and VRFs (from the range specified in fabric creation).
- Provision the overlay networks and VRFs on the switches of the fabric.
- Undeploy the networks and VRFs from the switches.
- Remove the provisioning from the fabric in Nexus Dashboard.

Manage > Inventory > Switches > Switch Overview

Provisioning of configuration on service leafs to which L4-7 service appliances may be attached.

This chapter mostly covers configuration provisioning for a single VXLAN BGP EVPN fabric. EVPN Multi-Site provisioning for Layer-2/Layer-3 DCI across multiple fabrics using the VXLAN fabric group, is documented in a separate chapter. The deployment details of how overlay Networks and VRFs can be easily provisioned from the Nexus Dashboard is covered in the "Networks" and "VRFs" sections in [Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric](#).

Guidelines for VXLAN BGP EVPN fabrics provisioning

- For any switch to be successfully imported into Nexus Dashboard, the user specified for discovery/import, should have the following permissions:
 - SSH access to the switch
 - Ability to perform SNMPv3 queries
 - Ability to run the **show** commands including show run, show interfaces, etc.
 - Ability to execute the **guestshell** commands, which are prefixed by **run guestshell** for the Nexus Dashboard tracker.
- The switch discovery user need not have the ability to make any configuration changes on the switches. It is primarily used for read access.
- When an invalid command is deployed by Nexus Dashboard to a device, for example, a command with an invalid key chain due to an invalid entry in the fabric settings, an error is generated displaying this issue. This error is not cleared after correcting the invalid fabric entry. You need to manually clean up or delete the invalid commands to clear the error.

Note that the fabric errors related to the command execution are automatically cleared only when the same failed command succeeds in the subsequent deployment.

- LAN credentials are required to be set of any user that needs to be perform any write access to the device. LAN credentials need to be set on the Nexus Dashboard, on a per user per device basis. When a user imports a device into the Easy Fabric, and LAN credentials are not set for that device, Nexus Dashboard moves this device to a migration mode. Once the user sets the appropriate LAN credentials for that device, a subsequent Save & Deploy retriggers the device

import process.

- The **Save & Deploy** button triggers the intent regeneration for the entire fabric as well as a configuration compliance check for all the switches within the fabric. This button is required but not limited to the following cases:
 - A switch or a link is added, or any change in the topology
 - A change in the fabric settings that must be shared across the fabric
 - A switch is removed or deleted
 - A new vPC pairing or unpairing is done
 - A change in the role for a device

When you click **Recalculate Config**, the changes in the fabric are evaluated, and the configuration for the entire fabric is generated. Click **Preview Config** to preview the generated configuration, and then deploy it at a fabric level. Therefore, **Deploy Config** can take more time depending on the size of the fabric.

+ When you right-click on a switch icon, you can use the **Deploy config to switches** option to deploy per switch configurations. This option is a local operation for a switch, that is, the expected configuration or intent for a switch is evaluated against its current running configuration, and a config compliance check is performed for the switch to get the **In-Sync** or **Out-of-Sync** status. If the switch is out of sync, the user is provided with a preview of all the configurations running in that particular switch that vary from the intent defined by the user for that respective switch.

- Persistent configuration diff is seen for the command line: `system nve infra-vlan int force`. The persistent diff occurs if you have deployed this command via the freeform configuration to the switch. Although the switch requires the force keyword during deployment, the running configuration that is obtained from the switch in Nexus Dashboard doesn't display the force keyword. Therefore, the `system nve infra-vlan int force` command always shows up as a diff.

The intent in Nexus Dashboard contains the line:

```
system nve infra-vlan int force
```

The running config contains the line:

```
system nve infra-vlan [int]
```

As a workaround to fix the persistent diff, edit the freeform config to remove the force keyword after the first deployment such that it is `system nve infra-vlan int`.

The force keyword is required for the initial deploy and must be removed after a successful deploy. You can confirm the diff by using the **Side-by-side Comparison** tab in the **Config Preview** window.

The persistent diff is also seen after a write erase and reload of a switch. Update the intent on Nexus Dashboard to include the force keyword, and then you need to remove the force keyword after the first deployment.

- When the switch contains the **hardware access-list tcam region arp-ether 256** command, which is deprecated without the **double-wide** keyword, the below warning is displayed:



Configuring the arp-ether region without "double-wide" is deprecated and can result in silent non-vxlan packet drops. Use the "double-wide" keyword when carving TCAM space for the arp-ether region.

Since the original **hardware access-list tcam region arp-ether 256** command doesn't match the policies in Nexus Dashboard, this config is captured in the **switch_freeform** policy. After the **hardware access-list tcam region arp-ether 256 double-wide** command is pushed to the switch, the original **tcam** command that does not contain the **double-wide** keyword is removed.

You must manually remove the **hardware access-list tcam region arp-ether 256** command from the **switch_freeform** policy. Otherwise, config compliance shows a persistent diff.

Here is an example of the **hardware access-list** command on the switch:

```
switch(config)# show run | inc arp-ether
switch(config)# hardware access-list tcam region arp-ether 256
Warning: Please save config and reload the system for the configuration to take effect
switch(config)# show run | inc arp-ether
hardware access-list tcam region arp-ether 256
switch(config)#
switch(config)# hardware access-list tcam region arp-ether 256 double-wide
Warning: Please save config and reload the system for the configuration to take effect
switch(config)# show run | inc arp-ether
hardware access-list tcam region arp-ether 256 double-wide
```

You can see that the original **tcam** command is overwritten.

AI fabric management

Nexus Dashboard 4.3.1 AI fabric settings introduces several enhancements, simplifying configuration and providing more granular control over network behavior. These updates streamline deployment and optimize performance for demanding AI workloads.

- Automated BGP ASN allocation: Automatic allocation of BGP Autonomous System Numbers (ASNs) for leaf switches, border devices, and border gateways in Multi-AS mode. This feature, enabled by default for new fabrics, assigns unique ASNs from a pre-defined pool, ensuring consistency for vPC switches.



You cannot simultaneously enable ASN auto allocation and the "Allow Same ASN on Leafs" flag.

- Increased BGP maximum paths: The default BGP maximum path value for all eBGP fabric types is now 64 (previously 4). This enhancement improves load balancing and network utilization.
- Default Routing Protocol for AI VXLAN EVPN Fabric: When creating a new AI VXLAN EVPN fabric, the default underlay and overlay routing protocol is now eBGP, aligning with common deployment

practices for AI.

- Revised default interface types: For greenfield AI fabric deployments, default interface types for host ports are revised:
 - AI BGP fabrics: Leaf host ports default to routed.
 - AI VXLAN EVPN eBGP/iBGP fabrics: Leaf host ports and ToR switches default to access.
- All AI-specific parameters are now consolidated under **AI settings** within the fabric settings. This simplifies access and management of critical configurations.
- Enhanced QoS Configuration for ECN and PFC: Nexus Dashboard exposes granular QoS parameters directly within the fabric settings, enabling automated enablement and threshold adjustments for Explicit Congestion Notification (ECN) and Priority Flow Control (PFC). This integrated approach includes:
 - PFC Enablement: Configurable parameters for CoS/PCP or DSCP for lossless handling, and a watchdog interval (default 100ms, valid range 101–1000ms).
 - ECN Enablement: Configurable parameters for WRED for buffer utilization, including WRED max/min thresholds and qualifications for flow differentiation.
 - Customization: You can provide DSCP values for QoS classification at the fabric level and configure DLB modes for specific traffic classes (if global DLB is policy-driven). Advanced users can match traffic by DSCP value and access list to select the DLB method.

AI settings

The **AI settings** centralizes critical parameters for configuring and managing AI network fabrics. This addresses the challenge of managing complex, specialized configurations by providing granular control over traffic management and hardware interactions for AI networks.

1. Enter the necessary information for the AI QoS & queuing policy.

Field	Description
AI QoS & queuing policy	Available for AI fabric types. Choose the queuing policy from the Policy selection mode drop-down list based on the predominant fabric link speed for certain switches in the fabric. For more information, see AI QoS classification and queuing policies. Options are: • 800G: Enable QoS queuing policies for an interface speed of 400 Gb. • 400G: Enable QoS queuing policies for an interface speed of 400 Gb. • 100G: Enable QoS queuing policies for an interface speed of 100 Gb. • 25G: Enable QoS queuing policies for an interface speed of 25 Gb. • User-defined: Provides customization options for parameters that are fixed in port speed templates.
	The rows below appear if you choose User-defined in the AI QoS & queuing policy field.
ROCEv2 DSCP	Enter the DSCP value (0–63) for RDMA traffic. The default value is 26.

Field	Description
CNP DSCP	Enter the DSCP value (0–63) for congestion notification. The default value is 48.
WRED minimum threshold (in kbytes)	Enter the minimum queue size for WRED. The default value is 950.
WRED maximum threshold (in kbytes)	Enter the maximum queue size for WRED. The default value is 3000.
Drop probability %	Enter the value to control WRED. The default value is 7.
WRED weight	Enter the value to control WRED, which influences queue changes. The default value is 0.
Bandwidth remaining %	Enter the percentage of remaining bandwidth allocated to AI traffic queue. The default value is 50.

2. Enter the necessary information for dynamic load balancing and job monitoring telemetry.

Field	Description
Dynamic load balancing	Click to Enable the options. Choose the Mode selection from the drop-down list and enter the aging period in the Flowlet aging time field. For more information, see Add a Dynamic Load Balancing (DLB) policy template.  - Flowlet aging is applicable when you choose the flowlet, policy driven flowlet, or policy driven mixed mode mode option from the Dynamic_Load_Balancing_mode drop-down list when you add a Dynamic Load Balancing (DLB) policy template. It is not effective when you choose the per packet or policy driven per packet mode option from the Dynamic_Load_Balancing_mode drop-down list. - The NX-OS CLI does accept a flowlet aging value, regardless of the mode; however, it may have a no operation result based on the mode choice.
Job monitoring telemetry	Displays the Server data that contains server related information such as the model and other types. The Update version allows you to upload the latest version of the zip file that you created with the GPU/NIC Topology Fleet Collector toolkit. See GPU/NIC Topology Fleet Collector for more information.
Delete mapping file	If you want to delete a GPU/NIC Topology Fleet Collector zip file that you uploaded previously, toggle the switch to Enable and save.

Field	Description
Telemetry data pull mode	Choose the telemetry data pull mode that you want to use. • Secured SSL: You will need a certificate to connect to the DCGM Exporter if you choose the Secured SSL telemetry data pull mode. See DCGM Exporter for more information. Enter the necessary information in the Client certification name field below, which appears when you choose the Secured SSL telemetry data pull mode. • Unsecured: This telemetry data pull mode is over HTTP. You do not have to enter any additional information if you choose the Unsecured telemetry data pull mode.
Client certification name	Enter the name of the client certification that you created using the procedures provided in Collect data in secured mode .

Layer 3 VNI without VLAN

Following is the upper-level process to enable the Layer 3 VNI without VLAN feature in a fabric:

1. (Optional) When configuring a new fabric, check the **Enable L3VNI w/o VLAN** field to enable the Layer 3 VNI without VLAN feature at the fabric level. The setting at this fabric-level field affects the related field at the VRF level, as described below.
2. When creating or editing a VRF, check the **Enable L3VNI w/o VLAN** field to enable the Layer 3 VNI without VLAN feature at the VRF level. The default setting for this field varies depending on the following factors:
 - For existing VRFs, the default setting is disabled (the **Enable L3VNI w/o VLAN** box is unchecked).
 - For newly-created VRFs, the default setting is inherited from the fabric settings, as described above.
 - This field is a per-VXLAN fabric variable. For VRFs that are created from a VXLAN EVPN Multi-Site fabric, the value of this field is inherited from the fabric setting in the child fabric. You can edit the VRF in the child fabric to change the value, if desired.

See the "Create a VRF" section in [Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric](#) for more information.

The VRF attachment (new or edited) then uses the new Layer 3 VNI without VLAN mode if the following conditions are met:

- The **Enable L3VNI w/o VLAN** is enabled at the VRF level
- The switch supports this feature and the switch is running on the correct release (see [Guidelines and limitations for Layer 3 VNI without VLAN](#))

The VLAN is ignored in the VRF attachment when these conditions are met.

Guidelines and limitations for Layer 3 VNI without VLAN

Following are the guidelines and limitations for the Layer 3 without VLAN feature:

- The Layer 3 VNI without VLAN feature is supported on the -EX, -FX, and -GX versions of the Cisco N9000 switches. When you enable this feature at the VRF level, the feature setting on the VRF will be ignored on switch models that do not support this feature.
- When used in a Campus VXLAN EVPN fabric, this feature is only supported on Cisco N9000 series switches in that type of fabric. This feature is not supported on Cisco Catalyst 9000 series switches in the Campus VXLAN EVPN fabric; those switches require VLANs for Layer 3 VNI configurations.
- This feature is supported on switches running on NX-OS release 10.3.1 or later. If you enable this feature at the VRF level, the feature setting on the VRF is ignored on switches running an NX-OS image earlier than 10.3.1.
- When you perform a brownfield import in a Data Center VXLAN EVPN fabric, if one switch configuration is set with the **Enable L3VNI w/o VLAN** configuration at the VRF level, then you should also configure this same setting for the rest of the switches in the same fabric that are associated with this VRF, if the switch models and images support this feature.
- For VXLAN fabric or fabric group deployments with Campus switches, use a VNI range within 4096 to 1,677,721 to align with the best practices for VNI allocation in Campus environments.

Precision Time Protocol for Data Center VXLAN EVPN fabrics

In the fabric settings for the **Data Center VXLAN EVPN** template, select the **Enable Precision Time Protocol (PTP)** check box to enable PTP across a fabric. When you select this check box, PTP is enabled globally and on core-facing interfaces. Additionally, the **PTP Loopback Id** and **PTP Domain Id** fields are editable.

The PTP feature works only when all the devices in a fabric are cloud-scale devices. Warnings are displayed if there are non-cloud scale devices in the fabric, and PTP is not enabled. Examples of the cloud-scale devices are Cisco N93180YC-EX, Cisco N93180YC-FX, Cisco N93240YC-FX2, and Cisco N93360YC-FX2 switches.

For more information, see the *Configuring PTP* chapter in *Cisco N9000 Series NX-OS System Management Configuration Guide* and *Cisco Nexus Dashboard Insights User Guide*.

For Nexus Dashboard Fabric Controller deployments, specifically in a VXLAN EVPN based fabric deployments, you have to enable PTP globally, and also enable PTP on core-facing interfaces. The interfaces could be configured to the external PTP server like a VM or Linux-based machine. Therefore, the interface should be edited to have a connection with the grandmaster clock.

We recommend that you configure the grandmaster clock outside of the fabric and that it is IP reachable. The interfaces toward the grandmaster clock need to be enabled with PTP via the interface freeform config.

All core-facing interfaces are auto-enabled with the PTP configuration after you click **Deploy Config**. This action ensures that all devices are PTP synced to the grandmaster clock. Additionally, for any interfaces that are not core-facing, such as interfaces on the border devices and leafs that are connected to hosts, firewalls, service-nodes, or other routers, the TTAG related CLI must be added.

The TTAG is added for all traffic entering the VXLAN EVPN fabric and the TTAG must be stripped when traffic is exiting this fabric.

Here is the sample PTP configuration:

```
feature ptp

ptp source 100.100.100.10 -> _IP address of the loopback interface (loopback0) that is
already created or user created loopback interface in the fabric settings_

ptp domain 1 -> _PTP domain ID specified in fabric settings_

interface Ethernet1/59 -> _Core facing interface_
 ptp

interface Ethernet1/50 -> _Host facing interface_
 ttag
 ttag-strip
```

The following guidelines are applicable for PTP:

- The PTP feature can be enabled in a fabric when all the switches in the fabric have Cisco NX-OS Release 7.0(3)I7(1) or a higher version. Otherwise, the following error message is displayed:

PTP feature can be enabled in the fabric, when all the switches have NX-OS Release 7.0(3)I7(1) or higher version. Please upgrade switches to NX-OS Release 7.0(3)I7(1) or higher version to enable PTP in this fabric.

- For hardware telemetry support in NIR, the PTP configuration is a prerequisite.
- If you are adding a non-cloud scale device to an existing fabric which contains PTP configuration, the following warning is displayed:

TTAG is enabled fabric wide, when all devices are cloud scale switches so it cannot be enabled for newly added non cloud scale device(s).

- If a fabric contains both cloud scale and non-cloud scale devices, the following warning is displayed when you try to enable PTP:

TTAG is enabled fabric wide, when all devices are cloud scale switches and is not enabled due to non cloud scale device(s).

AI QoS classification and queuing policies

These sections provide information about the AI QoS classification and queuing policies.

- [Understanding AI QoS classification and queuing policies](#)
- [Guidelines and limitations for AI QoS classification and queuing policies](#)
- [Configure AI QoS classification and queuing policies](#)

- [Create a policy using the custom QoS templates](#)

Understanding AI QoS classification and queuing policies

Support is available for configuring a low latency, high throughput, and lossless fabric configuration that can be used for artificial intelligence (AI) fabric based traffic.

The AI QoS feature allows you to:

- Easily configure a network with homogeneous interface speeds, where most or all of the links run at 400Gb, 100Gb, or 25Gb speeds.
- Provide customizations to override the predominate queuing policy for a host interface.

When you apply the AI QoS policy, Nexus Dashboard will automatically pre-configure any inter-fabric links with QoS and system queuing policies, and will also enable Priority Flow Control (PFC). If you enable the AI QoS feature on a VXLAN EVPN fabric, then the Network Virtual (NVE) interface will have the attached AI QoS policies.

You can enable this feature and set queuing policy parameters based on interface speed using new fields available during BGP fabric configuration.

+ Policies defined with these custom Classification and Queuing templates can be used in various host interface policies. For more information, see [Create a policy using the custom QoS templates](#).

When enabling the AI feature, **priority-flow-control watchdog-interval on** is enabled on all of your configured devices, intra-fabric links, and all your host interfaces where Priority Flow Control (PFC) is also enabled. The PFC watchdog interval is for detecting whether packets in a no-drop queue are being drained within a specified time period. This release also adds the **Priority flow control watch-dog interval** field on the **Advanced** tab. When you create or edit a Data Center VXLAN EVPN fabric or other fabrics and AI is enabled, you can set the **Priority flow control watch-dog interval** field to a non-system default value (the default is 100 milliseconds). For more information on the PFC watchdog interval for Cisco NX-OS, see [Configuring a priority flow control watchdog Interval](#) in the *Cisco N9000 Series NX-OS Quality of Service Configuration Guide*.

If you perform an upgrade from an earlier release, and then do a **Recalculate and deploy**, you may see additional **priority-flow-control watchdog-interval on** configurations.

Guidelines and limitations for AI QoS classification and queuing policies

Following are the guidelines and limitations for the AI QoS and queuing policy feature:

- Apply AI QoS policies at the fabric level rather than on individual switches to ensure consistent traffic classification and uniform policy enforcement.
- On Cisco N9K-C9808 and N9K-C9804 series switches, the command **priority-flow-control watch-dog-interval** is not supported in either global or interface configuration modes and the command **hardware qos nodrop-queue-thresholds queue-green** is not supported in global configuration mode.
- Cisco N9K-C9808 and N9K-C9804 series switches only support AI fabric type from NX-OS version 10.5(1) and later.
- This feature does not automate any per-interface speed settings.

- This feature is supported only on Nexus devices with Cisco Cloud Scale technology, such as the Cisco N9300-FX2, Cisco N9300-FX3, Cisco N9300-GX, and Cisco N9300-GX2 series switches.
- This feature is not supported in fabrics with devices that are assigned with a ToR role.

Configure AI QoS classification and queuing policies

Follow these steps to configure AI QoS and queuing policies:

1. Enable AI QoS and queuing policies at the fabric level.
 - a. Create a fabric as you normally would.
 - b. In the **Advanced** tab in those instructions, make the necessary selections to configure AI QoS and queuing policies at the fabric level.
 - c. Configure any remaining fabric-level settings as necessary in the remaining tabs.
 - d. When you have completed all the necessary fabric-level configurations, click **Save**, then click **Recalculate and deploy**.

At this point in the process, the network QoS and queuing policies are configured on each device, the classification policy is configured on NVE interfaces (if applicable), and priority flow control and classification policy is configured on all intra-fabric link interfaces.

2. For host interfaces, selectively enable priority flow control, QoS, and queuing by editing the policy associated with that host interface.

See [Working with Connectivity for LAN Fabrics](#) for more information.

- a. Within a fabric where you enabled AI QoS and queuing policies in the previous step, click the **Interfaces** tab.

The configured interfaces within this fabric are displayed.

- b. Locate the host interface where you want to enable AI QoS and queuing policies, then click the box next to that host interface to select it and click **Actions > Edit**.

The **Edit Interfaces** page is displayed.

- c. In the **Policy** field, verify that the policy that is associated with this interface contains the necessary fields that will allow you to enable AI QoS and queuing policies on this host interface.

For example, these policy templates contain the necessary AI QoS and queuing policies fields:

- int_access_host
- int_dot1q_tunnel_host
- int_pvlan_host
- int_routed_host
- int_trunk_host

- d. Locate the **Enable priority flow control** field and click the box next to this field to enable Priority Flow Control for this host interface.

- e. In the **Enable QoS Configuration** field, click the box next to this field to enable AI QoS for this host interface.

This enables the QoS classification on this interface if AI queuing is enabled at the fabric level.

- f. If you checked the box next to the **Enable QoS Configuration** field in the previous step and you created a custom QoS policy using the procedures provided in [Create a policy using the custom QoS templates](#), enter that custom QoS classification policy in the **Custom QoS Policy for this interface** field to associate that custom QoS policy with this host interface, if necessary.

If this field is left blank, then Nexus Dashboard will use the default QOS_CLASSIFICATION policy, if available.

- g. If you created a custom queuing policy using the procedures provided in [Create a policy using the custom QoS templates](#), enter that custom queuing policy in the **Custom Queuing Policy for this interface** field to associate that custom queuing policy with this host interface, if desired.
- h. Click **Save** when you have completed the AI QoS and queuing policy configurations for this host interface.

Create a policy using the custom QoS templates

Follow these procedures to use the custom QoS templates to create a policy, if desired. See [Managing Your Template Library](#) for general information on templates.

1. Within a fabric where you enabled AI QoS and queuing policies, click **Inventory > Switches**, then double-click the switch that has the host interface where you enabled AI QoS and queuing policies.

The **Switch overview** page for that switch appears.

2. Choose **Configuration Policies > Policies**.
3. Click **Actions > Add policy**.

The **Create Policy** page appears.

4. Set the priority and enter a description for the new policy.

Note that the priority for this policy must be lower (must come before) the priority that was set for the host interface.

5. In the **Select Template** field, click the **No Policy Selected** text.

The **Select Policy Template** page appears.

1. Make the necessary QoS classification or queuing configurations in the template that you selected, then click **Save**.

Any custom QoS policy created using these procedures are now available to use when you configure QoS and queuing policies for the host interface.

MACsec support in Data Center VXLAN EVPN and BGP fabrics

MACsec is supported in Data Center VXLAN EVPN and BGP fabrics for intra-fabric links. You should enable MACsec on the fabric and on each required intra-fabric link to configure MACsec. Unlike CloudSec, auto-configuration of MACsec is not supported.

Support is available for MACsec for inter-fabric links, DCI MACsec, with a QKD server to generate keys or using preshared keys. For more information, see the section "About connecting two NX-OS fabrics with MACsec using QKD" in [Creating Fabrics and Fabric Groups](#).

MACsec is supported on switches with minimum Cisco NX-OS releases 7.0(3), I7(8), and 9.3(5).

Guidelines

- If MACsec cannot be configured on the physical interfaces of the link, an error is displayed when you click **Save**. MACsec cannot be configured on the device and link due to the following reasons:
 - The minimum NX-OS version is not met.
 - The interface is not MACsec capable.
- MACsec global parameters in the fabric settings can be changed at any time.
- MACsec and CloudSec can coexist on a 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 configs.

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

The following sections show how to enable and disable MACsec in Nexus Dashboard.

Enable MACsec

The process described in this section is for configuring MACsec for intra-fabric links and inter-fabric links.

For information on configuring data center interconnect (DCI) MACsec for inter-fabric links and using a quantum key distribution (QKD) server, see the section "About connecting two NX-OS fabrics with MACsec using QKD" in [Creating Fabrics and Fabric Groups](#).

1. Navigate to **Manage > Fabrics**.
2. Click **Create Fabric** to create a new fabric or click **Actions > Edit Fabric Settings** on an existing Data Center VXLAN EVPN fabric.
3. Click **Fabric Management > Security** and specify the MACsec details.

Field	Description
Enable MACsec	Check the check box to enable MACsec in the fabric. MACsec configuration is not generated until MACsec is enabled on an intra-fabric link. Perform a Recalculate and deploy operation to generate the MACsec configuration and deploy the configuration on the switch.
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.
MACsec Primary Key String	Specify a Cisco Type 7 encrypted octet string that is used for establishing the primary 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.
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 .
Enable DCI MACsec	Check the check box to enable DCI MACsec on DCI links.  If you enable the Enable DCI MACsec option and disable the Use Link MACsec Setting option, Nexus Dashboard uses the fabric settings for configuring DCI MACsec on the DCI links.

Field	Description
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 generates preshared keys 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-XPB-128 • GCM-AES-XPB-256 The default value is GCM-AES-XPB-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.  If you disabled the Enable QKD option, you need to specify the DCI MACsec Fallback Key String option.
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.
QKD Profile Name	Specify the crypto profile name. The maximum size is 63.

Field	Description
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**.
- Click a fabric to view the **Summary** in the side panel.
- Click the **Launch** icon to open the **Fabric Overview** page.
- Click the **Fabric Overview > Links** tab.
- Choose an inter-fabric connection (IFC) link on which you want to enable MACsec and click **Actions > Edit**. For more information on creating a VRF-Lite inter-fabric link, see the section "Establishing inter-fabric connectivity using VRF Lite" [Working with Connectivity in Your Nexus Dashboard LAN Fabrics](#).

The **Link Management - Edit Link** page displays.

- From the **Link Management - Edit Link** page, navigate to the **Security** tab and enter the required parameters.

For an inter-fabric link, the **Enable MACsec** option is in [Security](#).

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 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 key-chain names for the MACsec configuration at the source. The default value is DCI, and you can change the value.

Field	Description
Destination MACsec Policy/Key-Chain Name Prefix	Specify the prefix for the policy and key-chain names for the MACsec configuration at the destination. The default value is DCI, and you can change the value.
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.
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-XPN-128 • GCM-AES-XPN-256 The default value is GCM-AES-XPN-256.
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.

Field	Description
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 KMEserver.
Source Trustpoint Label	Specify the source authentication type 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.

10. Click **Save**.

11. From the **Fabric Overview > Switches** tab, select **Actions > Recalculate and deploy** to deploy the MACsec configuration on the switch.

Disable MACsec

The process described in this section is for disabling MACsec on an inter-fabric link for Data Center VXLAN EVPN and BGP fabrics.

For information on disabling MACsec for an inter-fabric link with a quantum key distribution (QKD) server, see the section "About connecting two NX-OS fabrics with MACsec using QKD" in [Creating Fabrics and Fabric Groups](#)].

Nexus Dashboard performs the following operations when you disable MACsec:

- Disables MACsec on an 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 **Fabric Overview > Links** tab.
 2. Choose the link on which you want to disable MACsec on the inter-fabric link.

3. From the **Link Management - Edit Link** page, navigate to the appropriate tab and unselect the **Enable MACsec** option.

For an intra-fabric link, the **Enable MACsec** option is in the [Advanced](#) tab.

For an inter-fabric link, the **Enable MACsec** option is in the [Security](#) tab.

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

Provisioning VXLAN EVPN Fabric with IGP underlay

A fabric workflow is available for unified underlay and overlay provisioning of VXLAN EVPN configuration on Cisco N9000 and Cisco Nexus 3000 Series switches. The configuration of the fabric is achieved via a powerful, flexible, and customizable template-based framework. Using minimal user inputs, you can bring up the entire fabric with Cisco recommended best practice configurations, in a short period of time. The set of parameters exposed in the Fabric Settings allows you to tailor the fabric to their preferred underlay provisioning options.

For creating and deploying VXLAN EVPN fabrics, see [Creating Fabrics and Fabric Groups](#).

Create a VXLAN EVPN fabric with IPv4 underlay

To create a new VXLAN EVPN fabric, refer to [Creating Fabrics and Fabric Groups](#).

Create a VXLAN EVPN fabric with IPv6 underlay

This procedure shows how to create a VXLAN EVPN fabric with an IPv6 underlay. Note that only the fields for creating a VXLAN fabric with an IPv6 underlay are documented. For information about the remaining fields, see [Creating Fabrics and Fabric Groups](#).

Nexus Dashboard added support for configuring a Data Center VXLAN EVPN fabric and a BGP VXLAN fabric with an IPv6 PIMv6 underlay. For more information, see [Configuring VXLANv6 with a PIMv6 underlay and TRMv6 for a Data Center VXLAN EVPN fabric](#).

1. Choose **Manage > Fabrics**.
2. From the **Actions** drop-down list, choose **Create Fabric**.

The **Create Fabric** page appears.

3. Enter the name of the fabric in the **Fabric Name** field.
4. In the **Pick Fabric** field, choose **Data Center VXLAN EVPN** from the **Select Type of Fabric** drop-down list.
5. Click **Select**.
6. The **General Parameters** tab is displayed by default. The fields for configuring an IPv6 underlay in the **General Parameters** tab are the following:

Field	Description
BGP ASN	Enter the BGP AS number for the fabric. You can enter either the 2 byte BGP ASN or 4 byte BGP ASN.
Enable IPv6 Underlay	Check the Enable IPv6 Underlay check box.
Enable IPv6 Link-Local Address	Check the Enable IPv6 Link-Local Address check box to use the link local addresses in the fabric between leaf-spine and spine-border interfaces. If you check this check box, the Underlay Subnet IPv6 Mask field is not editable. By default, the Enable IPv6 Link-Local Address field is enabled.
Fabric Interface Numbering	An IPv6 underlay supports p2p networks only. Therefore, the Fabric Interface Numbering drop-down list is disabled.
Underlay Subnet IPv6 Mask	Specify the subnet mask for the fabric interface for IPv6 addresses.
Underlay Routing Protocol	Specify the IGP used in the fabric, that is, OSPF or IS-IS for VXLANv6.
Enable IPv6 Underlay	Check the Enable IPv6 Underlay check box.
Enable IPv6 Link-Local Address	Check the Enable IPv6 Link-Local Address check box to use the link local addresses in the fabric between leaf-spine and spine-border interfaces. If you check this check box, the Underlay Subnet IPv6 Mask field is not editable. By default, the Enable IPv6 Link-Local Address field is enabled. IPv6 underlay supports p2p networks only. Therefore, the Fabric Interface Numbering drop-down list is disabled.
Underlay Subnet IPv6 Mask	Specify the subnet mask for the fabric interface for IPv6 addresses.
Underlay Routing Protocol	Specify the IGP used in the fabric, that is, OSPF or IS-IS for VXLANv6.

- Navigate to the **Replication** tab. The fields for configuring an IPv6 underlay for the Multicast replication mode are the following:

Field	Description
Replication Mode	The mode of replication that is used in the fabric for BUM (Broadcast, Unknown Unicast, Multicast) traffic. The choices are Ingress replication or Multicast replication. When you choose Ingress replication, the multicast-related fields are disabled. When you choose Multicast replication, the Ingress replication fields are disabled. You can change the fabric setting from one mode to the other, if no overlay profile exists for the fabric. Choose Multicast as the replication mode for Tenant Routed Multicast (TRM) for IPv4 or IPv6. For more information for the IPv4 use case, see Editing AI Data Center VXLAN fabric settings. For more information for the IPv6 use case, see Configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6.
IPv6 Multicast Group Subnet	Enter an IPv6 multicast address with a prefix range of 112 to 126.
Enable IPv4 Tenant Routed Multicast (TRM)	Check this check box to enable Tenant Routed Multicast (TRM) with IPv4 that allows overlay IPv4 multicast traffic to be supported over EVPN/MVPN in VXLAN EVPN fabrics.
Enable IPv6 Tenant Routed Multicast (TRM)	Check the check box to enable Tenant Routed Multicast (TRM) with IPv6 that allows overlay IPv6 multicast traffic to be supported over EVPN/MVPN in VXLAN EVPN fabrics.
Default MDT IPv4 Address for TRM VRFs	The multicast address for Tenant Routed Multicast traffic is populated. By default, this address is from the IP prefix specified in the Multicast Group Subnet field. When you update either field, ensure that the address is chosen from the IP prefix specified in Multicast Group Subnet. For more information, see Overview of Tenant Routed Multicast.
Default MDT IPv6 Address for TRM VRFs	The multicast address for Tenant Routed Multicast traffic is populated. By default, this address is from the IP prefix specified in the IPv6 Multicast Group Subnet field. When you update either field, ensure that the address is chosen from the IP prefix specified in IPv6 Multicast Group Subnet. For more information, see Overview of Tenant Routed Multicast.

Field	Description
MVPN VRI ID Range	Use this field for allocating a unique MVPN VRI ID per vPC. This field is needed for the following use cases: ▪ If you configure a VXLANv4 underlay with a Layer 3 VNI without VLAN mode, and you enable TRMv4 or TRMv6. ▪ If you configure a VXLANv6 underlay, and you enable TRMv4 or TRMv6.  The MVPN VRI ID cannot be the same as any site ID within a multi-site fabric. The VRI ID has to be unique within all sites within an MSD.
Enable MVPN VRI ID Re-allocation	Enable this check box to generate a one-time VRI ID reallocation. Nexus Dashboard automatically allocates a new MVPN ID within the MVPN VRI ID range above for each applicable switch. Since this is a one-time operation, after performing the operation, this field is turned off.  Changing the VRI ID is disruptive, so plan accordingly.

8. Navigate to the **VPC** tab.

Field	Description
vPC Peer Keep Alive option	Choose management or loopback. To use IP addresses assigned to the management port and the management VRF, choose management. To use IP addresses assigned to loopback interfaces and a non-management VRF, choose underlay routing loopback with an IPv6 address for PKA. Both the options are supported for an IPv6 underlay.

9. Navigate to the **Protocols** tab.

Field	Description
Underlay Anycast Loopback Id	Specify the underlay anycast loopback ID for an IPv6 underlay. You cannot configure an IPv6 address as secondary, an additional loopback interface is allocated on each vPC device. Its IPv6 address is used as the VIP.

10. Navigate to the **Resources** tab.

Field	Description
Manual Underlay IP Address Allocation	Check this check box to manually allocate underlay IP addresses. The dynamic underlay IP addresses fields are disabled.
Underlay Routing Loopback IPv6 Range	Specify the loopback IPv6 addresses for protocol peering.
Underlay VTEP Loopback IPv6 Range	Specify the loopback IPv6 addresses for VTEPs.

Field	Description
Underlay Subnet IPv6 Range	Specify the IPv6 address range that is used for assigning IP addresses for numbered and peer link SVIs. To edit this field, uncheck the Enable IPv6 Link-Local Address check box under the General Parameters tab.
BGP Router ID Range for IPv6 Underlay	Specify the address range to assign BGP Router IDs. The IPv4 addressing is used for routers with BGP and underlay routing protocols.

11. Navigate to the **Bootstrap** tab.

Field	Description
Enable Bootstrap	Check the Enable Bootstrap check box. If this check box is not chosen, none of the other fields on this tab are editable.
Enable Local DHCP Server	Check the check box to initiate automatic assignment of IP addresses assignment through the local DHCP server. The DHCP Scope Start Address and the DHCP Scope End Address fields are editable only after you check this check box.
DHCP Version	Choose DHCPv4 from the drop-down list.

12. Navigate to the **Advanced** tab.

Field	Description
Allow L3VNI w/o VLAN	Check this check box to allow Layer 3 VNI configuration without having to configure a VLAN.
Enable TCAM Allocation	Check this option to provide TCAM allocation to 768 or above if you enable TRMv6. TCAM commands are automatically generated for VXLAN and vPC fabric peering when enabled. For more information, see Configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6.

13. Click **Save** to complete the creation of the fabric.

What's next: See the section "Add switches to a fabric" in [Working with Inventory in Your Nexus Dashboard LAN or IPFM Fabrics](#).

Add switches

Switches can be added to a single fabric at any point in time. To add switches to a fabric and discover existing or new switches, refer to the section "Add switches to a fabric" in [Working with Inventory in Your Nexus Dashboard LAN or IPFM Fabrics](#).

Beginning with Nexus Dashboard release 4.3.1, Nexus Dashboard adds support for Cisco N9348Y12C-SE1 switches running NX-OS 10.6(3) release on all LAN fabric types.

Assigning Switch Roles

To assign roles to switches on Nexus Dashboard refer to the section "Assign switch roles" in

vPC fabric peering

vPC Fabric Peering provides an enhanced dual-homing access solution without the overhead of wasting physical ports for vPC Peer Link. This feature preserves all the characteristics of a traditional vPC. For more information, see *Information about vPC Fabric Peering* section in *Cisco N9000 Series NX-OS VXLAN Configuration Guide*.

You can create a virtual peer link for two switches or change the existing physical peer link to a virtual peer link. Nexus Dashboard supports vPC fabric peering in both greenfield as well as brownfield deployments. This feature is applicable for **Data Center VXLAN EVPN** and **BGP Fabric** fabric templates.



The **BGP Fabric** fabric does not support brownfield import.

Guidelines and limitations

The following are the guidelines and limitations for vPC fabric pairing.

- vPC fabric peering is supported from Cisco NX-OS Release 9.2(3).
- Only the Cisco N9K-C9332C Switch, Cisco N9K-C9364C Switch, Cisco N9K-C9348GC-FXP Switch, and Cisco N9000 Series Switches that end with FX or FX2 support vPC fabric peering.
- Cisco N9K-C93180YC-FX3S and N9K-C93108TC-FX3P platform switches support vPC fabric peering.
- Cisco N9300-EX, and 9300-FX/FXP/FX2/FX3/GX/GX2 platform switches support vPC Fabric Peering. Cisco N9200 and Cisco N9500 platform switches do not support vPC Fabric Peering. For more information, see *Guidelines and Limitations for vPC Fabric Peering* section in *Cisco N9000 Series NX-OS VXLAN Configuration Guide*.
- If you use other Cisco N9000 Series Switches, a warning will appear during **Recalculate & Deploy**. A warning appears in this case because these switches will be supported in future releases.
- If you try pairing switches that do not support vPC fabric peering, using the **Use Virtual Peerlink** option, a warning will appear when you deploy the fabric.
- You can convert a physical peer link to a virtual peer link and vice-versa with or without overlays.
- Switches with border gateway leaf roles do not support vPC fabric peering.
- vPC fabric peering is not supported for Cisco N9000 Series Modular Chassis and FEXs. An error appears during **Recalculate & Deploy** if you try to pair any of these.
- Brownfield deployments and greenfield deployments support vPC fabric peering in Cisco Nexus Dashboard.
- However, you can import switches that are connected using physical peer links and convert the physical peer links to virtual peer links after **Recalculate & Deploy**. To update a TCAM region during the feature configuration, use the hardware access-list tcam ingress-flow redirect512 command in the configuration terminal.

QoS for fabric vPC-peering

In the **Data Center VXLAN EVPN** fabric settings, you can enable QoS on spines for guaranteed delivery of vPC Fabric Peering communication. Additionally, you can specify the QoS policy name.

Note the following guidelines for a greenfield deployment:

- If QoS is enabled and the fabric is newly created:
 - If spines or super spines neighbor is a virtual vPC, make sure neighbor is not honored from invalid links, for example, super spine to leaf or borders to spine when super spine is present.
 - Based on the Cisco N9000 Series Switch model, create the recommended global QoS config using the **switch_freeform** policy template.
 - Enable QoS on fabric links from spine to the correct neighbor.
- If the QoS policy name is edited, make sure policy name change is honored everywhere, that is, global and links.
- If QoS is disabled, delete all configuration related to QoS fabric vPC peering.
- If there is no change, then honor the existing PTI.

For more information about a greenfield deployment, see [Creating Fabrics and Fabric Groups](#).

Note the following guidelines for a brownfield deployment:

Brownfield Scenario 1:

- If QoS is enabled and the policy name is specified:



You need to enable only when the policy name for the global QoS and neighbor link service policy is same for all the fabric vPC peering connected spines.

- Capture the QoS configuration from switch based on the policy name and filter it from unaccounted configuration based on the policy name and put the configuration in the **switch_freeform** with PTI description.
- Create service policy configuration for the fabric interfaces as well.
- Greenfield configuration should make sure to honor the brownfield configuration.
- If the QoS policy name is edited, delete the existing policies and brownfield extra configuration as well, and follow the greenfield flow with the recommended configuration.
- If QoS is disabled, delete all the configuration related to QoS fabric vPC peering.



No cross check for possible or error mismatch user configuration, and user might see the diff.

Brownfield Scenario 2:

- If QoS is enabled and the policy name is not specified, QoS configuration is part of the unaccounted switch freeform config.
- If QoS is enabled from fabric settings after **Recalculate & Deploy** for brownfield, QoS configuration overlaps and you will see the diff if fabric vPC peering config is already present.

For more information about a brownfield deployment, see [Creating Fabrics and Fabric Groups](#).

To view the vPC pairing page of a switch, from the fabric topology page, right-click the switch and choose **vPC Pairing**. The vPC pairing page for a switch has the following fields:

Field	Description
Use Virtual Peerlink	Allows you to enable or disable the virtual peer linking between switches.
Switch name	Specifies all the peer switches in a fabric.NOTE: When you have not paired any peer switches, you can see all the switches in a fabric. After you pair a peer switch, you can see only the peer switch in the vPC pairing page.
Recommended	Specifies if the peer switch can be paired with the selected switch. Valid values are true and false . Recommended peer switches will be set to true .
Reason	Specifies why the vPC pairing between the selected switch and the peer switches is possible or not possible.
Serial Number	Specifies the serial number of the peer switches.

You can perform the following with the **vPC Pairing** option:

Create a virtual peer link

Follow these steps to create a virtual peer link.

1. Choose **Home > Topology**.
2. Choose a fabric with the **Data Center VXLAN EVPN** or **BGP Fabric** fabric type.
3. Right-click a switch and choose **vPC Pairing** from the drop-down list.

The page to choose the peer appears.



Alternatively, you can also navigate to the fabric **Overview** page. Choose a switch in the **Inventory > Switches** tab and click on **Actions > vPC Pairing** to create, edit, or unpair a vPC pair. However, you can use this option only when you choose a Cisco Nexus switch.

You will get the following error when you choose a switch with the border gateway leaf role.
<switch-name> has a Network/VRF attached. Please detach the Network/VRF before vPC Pairing/Unpairing

4. Check the **Use Virtual Peerlink** check box.
5. Choose a peer switch and check the **Recommended** column to see if pairing is possible.

If the value is **true**, pairing is possible. You can pair switches even if the recommendation is **false**. However, you will get a warning or error during **Recalculate & Deploy**.

6. Click **Save**.
7. In the **Topology** page, choose **Recalculate & Deploy**.

The **Deploy Configuration** page appears.

8. Click the field against the switch in the **Preview Config** column.

The **Config Preview** page appears for the switch.

9. View the vPC link details in the pending configuration and side-by-side configuration.
10. Close the page.
11. Click the pending errors icon next to **Recalculate & Deploy** icon to view errors and warnings, if any.

If you see any warnings that are related to TCAM, click the **Resolve** icon. A confirmation dialog box about reloading switches appears. Click **OK**. You can also reload the switches from the topology page. For more information, see *Guidelines and Limitations for vPC Fabric Peering and Migrating from vPC to vPC Fabric Peering* sections in *Cisco N9000 Series NX-OS VXLAN Configuration Guide*.

The switches that are connected through vPC fabric peering, are enclosed in a gray cloud.

Convert a physical peer link to a virtual peer link

Before you begin

- Perform the conversion from physical peer link to virtual peer link during the maintenance window of switches.
- Ensure the switches support vPC fabric peering. Only the following switches support vPC fabric peering:
 - Cisco N9K-C9332C Switch, Cisco N9K-C9364C Switch, and Cisco N9K-C9348GC-FXP Switch.
 - Cisco N9000 Series Switches that ends with FX, FX2, and FX2-Z.
 - Cisco N9300-EX, and Cisco N9300-FX/FXP/FX2/FX3/GX/GX2 platform switches. For more information, see *Guidelines and Limitations for vPC Fabric Peering* section in *Cisco N9000 Series NX-OS VXLAN Configuration Guide*.

Follow these steps to convert a physical peer link to a virtual peer link.

1. Choose **Home > Topology**.
2. Choose a fabric with the **Data Center VXLAN EVPN** or **BGP Fabric** fabric type.
3. Right-click the switch that is connected using the physical peer link and choose **vPC Pairing** from the drop-down list.

The page to choose the peer appears.



Alternatively, you can also navigate to the fabric **Overview** page. Choose a switch in the **Inventory > Switches** tab and click on **Actions > vPC Pairing** to create, edit, or unpair a vPC pair. However, you can use this option only when

you choose a Cisco Nexus switch.

You will get the following error when you choose a switch with the border gateway leaf role.
<switch-name> has a Network/VRF attached. Please detach the Network/VRF before vPC Pairing/Unpairing

4. Check the **Recommended** column to see if pairing is possible.

If the value is **true**, pairing is possible. You can pair switches even if the recommendation is **false**. However, you will get a warning or error during **Recalculate & Deploy**.

5. Check the **Use Virtual Peerlink** check box.

The **Unpair** icon changes to **Save**.

6. Click **Save**.



After you click **Save**, the physical vPC peer link is automatically deleted between the switches even without deployment.

7. In the **Topology** page, choose **Recalculate & Deploy**.

The **Deploy Configuration** page appears.

8. Click the field against the switch in the **Preview Config** column.

The **Config Preview** page appears for the switch.

9. View the vPC link details in the pending configuration and the side-by-side configuration.

10. Close the page.

11. Click the pending errors icon next to the **Recalculate & Deploy** icon to view errors and warnings, if any.

If you see any warnings that are related to TCAM, click the **Resolve** icon. A confirmation dialog box about reloading switches appears. Click **OK**. You can also reload the switches from the fabric topology page.

The physical peer link between the peer switches turns red. Delete this link. The switches are connected only through a virtual peer link and are enclosed in a gray cloud.

Convert a virtual peer link to a physical peer link

Before you begin

Connect the switches using a physical peer link before disabling the vPC fabric peering.

Follow these steps to convert a virtual peer link to a physical peer link.

1. Choose **Home > Topology**.
2. Choose a fabric with the **Data Center VXLAN EVPN** or **BGP Fabric** fabric type.
3. Right-click the switch that is connected through a virtual peer link and choose **vPC Pairing** from the drop-down list.

The page to choose the peer appears.



Alternatively, you can also navigate to the fabric **Overview** page. Choose a switch in the **Inventory > Switches** tab and click on **Actions > vPC Pairing** to create, edit, or unpair a vPC pair. However, you can use this option only when you choose a Cisco Nexus switch.

4. Uncheck the **Use Virtual Peerlink** check box.

The **Unpair** icon changes to **Save**.

5. Click **Save**.
6. In the **Topology** page, choose **Recalculate & Deploy**.

The **Deploy Configuration** page appears.

7. Click the field against the switch in the **Preview Config** column.

The **Config Preview** page appears for the switch.

8. View the vPC peer link details in the pending configuration and the side-by-side configuration.
9. Close the page.
10. Click the pending errors icon next to the **Recalculate & Deploy** icon to view errors and warnings, if any.

If you see any warnings that are related to TCAM, click the **Resolve** icon. The confirmation dialog box about reloading switches appears. Click **OK**. You can also reload the switches from the fabric topology page.

The virtual peer link, represented by a gray cloud, disappears and the peer switches are connected through a physical peer link.

Overlay mode

You can create a VRF or network in CLI or config-profile mode at the fabric level. The overlay mode of member fabrics of a VXLAN EVPN Multi-Site fabric is set individually at the member-fabric level. Overlay mode can only be changed before deploying overlay configurations to the switches. After the overlay configuration is deployed, you cannot change the mode unless all the VRF and network attachments are removed.

Beginning with Nexus Dashboard 4.3.1, you can migrate overlay mode of a fabric from config-profile to CLI, even when overlays with active attachments exist. This migration process enables a seamless transition without removing existing configurations, under these specific conditions.

- All switches in the fabric must be running NX-OS 10.5(5) or later.
- All switches must be in a synchronized state (in-sync) before initiating the migration.



If no overlays or overlay attachments exist, you can change the mode immediately without triggering the migration workflow.

Limitations

When you change the overlay mode from config-profile to CLI in a fabric with existing attachments, Nexus Dashboard initiates a migration workflow. During this process, Nexus Dashboard temporarily blocks most user operations to ensure data consistency.

Nexus Dashboard blocks these actions until the migration completes.

- All overlay operations per fabric
- MSD overlay definition add/edit
- MSD preview/deploy (if the selected switch/attachment's fabric is in migration mode)
- Periodic **copy running-config startup-config (copy r s)** operations
- Backup operations
- Fabric settings change except overlay mode change
- One manage overlay operation
- MSD add/remove fabric
- Interface group actions

Nexus Dashboard allows these actions during migration.

- Recalculate
- Preview
- Deploy
- Interface actions, including deploy and push config
- Fabric restore operations



When you try to save fabric settings, Nexus Dashboard generates a migration warning alarm. If you attempt to perform any of the blocked actions, Nexus Dashboard displays an error message. During migration, Nexus Dashboard ensures you do not lose overlay connectivity or intent. However, if the fabric or device does not support it, you cannot move switches to maintenance mode during migration.

Brownfield import

If you upgrade from Cisco DCNM Release 11.5(x) or perform a brownfield import, you can import switches with config-profile-based overlays only in config-profile overlay mode. If you try to import them in CLI overlay mode, Nexus Dashboard displays an error. After a brownfield import where overlays use config-profile mode, you can use the config-profile to CLI migration feature to convert the overlay mode to CLI (provided the preconditions are met).

Configure overlay mode

Follow these steps to configure the overlay mode of VRFs or networks in a fabric.

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

Go to **Manage > Fabrics**.

2. Choose the desired fabric.
3. From the **Actions** drop-down list, choose **Edit fabric settings**.

The **Edit *fabric_name* Settings** page opens.

4. In the **Edit *fabric_name* Settings** page, click the **Fabric management** tab.
5. Click the **Advanced** tab.
6. From the **Overlay Mode** drop-down list, choose one of the following:
 - **config-profile** – for using configuration profiles
 - **cli** – for direct command-line interface configuration
7. Click **Save**.

Configuring downstream VNI

In a VXLAN fabric, the Layer 3 and Layer 2 virtual network identifier (VNIs) are centrally managed using a VXLAN fabric group for inter-fabric connectivity. All the VXLAN fabrics within the VXLAN fabric group use the same VNI value for the Layer 3 or Layer 2 network. The problem is that before the fabrics are brought in for inter-fabric connectivity, the fabrics have been managed as standalone fabrics with existing VRFs and networks. The Layer 2 and Layer 3 overlays have been configured independently and have conflicting VNIs among fabrics.

There are two types of VNI conflicts:

- The same VNI is used by different VRFs or networks in different VXLAN fabrics.
- The same VRF or network uses different VNIs in different VXLAN fabrics.

This feature is supported when creating or editing fabric types:

- VXLAN
- Campus VXLAN Prior to Nexus Dashboard release 4.1.1, you could not add a VXLAN fabric to a VXLAN fabric group if there was a VNI conflict. With Nexus Dashboard release 4.1.1, downstream VNI (DSVNI) allows VXLAN fabrics with a VNI conflict to communicate with each other. On the border gateway, Nexus Dashboard uses different VNIs for exchanging intra-fabric and inter-fabric traffic. The existing VNI continues to be used for intra-fabric traffic. Stitching of the VNI occurs at the border gateways for inter-fabric traffic between fabrics using different VNIs.

Nexus Dashboard added downstream VNI options in a VXLAN fabric group for configuring a global Layer 2 VNI and a Layer 3 VNI pool. The two ranges should not conflict with VNIs already in use in existing VXLAN fabrics. We suggest picking from the high end of the VNI range for the global **Layer 3 VXLAN VNI Global Range** and the **Layer 2 VXLAN VNI Global Range** in the **General Parameters** page for the fabric. Nexus Dashboard generates a new VRF and a network VNI allocation based on these two ranges.



Enabling downstream VNI in a VXLAN fabric group does not affect existing VRFs and networks.

When Nexus Dashboard allocates a VNI for a VRF or a network, and its intra-fabric VNI is different from the downstream VNI, Nexus Dashboard requires additional configuration on the border gateway.

VRF CLIs on border gateways for downstream VNI

```
ip extcommunity-list standard <vrf-name> seq 10 permit rt 2324:50005
route-map MS-FABRIC-TO-EXTERNAL-RMAP permit 100
  match extcommunity <vrf-name>
  set extcomm-list <vrf-name> delete

vrf context <vrf-name>
  address-family ipv4 unicast
    route-target both <local-asn>:<DSVNI>
    route-target both <local-asn>:<DSVNI> evpn
  address-family ipv6 unicast
    route-target both <local-asn>:<DSVNI>
    route-target both <local-asn>:<DSVNI> evpn
```

Network CLIs on border gateways for downstream VNI

```
ip extcommunity-list standard <network-name> permit rt 27:30001
route-map MS-FABRIC-TO-EXTERNAL-RMAP permit 200
  match extcommunity <network-name>
  set extcomm-list <network-name> delete
route-map MS-FABRIC-TO-EXTERNAL-RMAP permit 65535
evpn
  vni <local-VNI> I2
  route-target both <local-asn>:<DSVNI>
```

Nexus Dashboard generates an additional route target containing the downstream VNI allowing border gateways in different fabrics to exchange routes.

extcommunity-list and **route-map** removes the local VNI from the BGP updates towards Data Center Interconnectivity (DCI). This prevents route leaking if the same VNI is used in a different fabric for a different VRF or network. **route-map MS-FABRIC-TO-EXTERNAL-RMAP** is applied to all multi-site overlay inter-fabric links if downstream VNI is enabled in the fabric group, regardless of whether the inter-fabric link is manually created or auto-generated by Nexus Dashboard.

Benefits of downstream VNI

- Resolves overlapping VNIs when using conflicted VNIs for a VRF or a network in different fabrics because of not changing the default VNI pool, which is the same for all VXLAN fabrics
- Supports route exchange using a route target
- Prevents route leaking
- The downstream VNI feature provides normalized VNI ID support for both Layer 2 and Layer 3 VNIs, helping manage VNI consistency and prevent conflicts.

Use cases for downstream VNI

When you add a new VXLAN fabric to a VXLAN fabric group for inter-fabric connectivity, and if Nexus Dashboard detects a VNI conflict, Nexus Dashboard allocates a new downstream VNI range.

- If the same VNI is used by a VRF or a network in a VXLAN fabric group and the incoming VXLAN fabric (**vrf1** in the fabric group and **vrf2** in the incoming VXLAN fabric), Nexus Dashboard allocates a new VNI from the global VNI pool for both the VRF and the network. In this example, Nexus Dashboard allocates a new VNI for **vrf1** and **vrf2**.
- If a VRF or a network use a different VNI in a VXLAN fabric group and the incoming VXLAN fabric, Nexus Dashboard allocates a new VNI, if the VNI in the VXLAN fabric group is not already in the global VNI range.
- In these two use cases, if there is a VNI conflict between the VXLAN fabric group and the incoming VXLAN fabric, and the VNI in the incoming VXLAN fabric is already in the global VNI range, you cannot add the VXLAN fabric to the VXLAN fabric group. In this use case, you can edit the global **L2 VNI** and or the **L3 VNI** range so that the VNI range does not overlap with the VNIs already used in the incoming VXLAN fabric.

Supported platforms

Ensure that all border gateways have the correct platform and NX-OS version that supports downstream VNI. For the list of supported platforms and NX-OS versions, see the [Cisco N9000 Series NX-OS VXLAN Configuration Guide](#).

Guidelines and limitations for downstream VNI

- You must use unique names for VRFS and networks.

This means that **vrf foo** on fabric1 and **vrf foo** on fabric2 refer to the same VRF. The same applies for network names.

- You cannot disable downstream VNI in a VXLAN fabric group if there are any existing VRFS or networks with a fabric VNI that differs from the VNI of the VXLAN fabric group.

These features are not supported for downstream VNI:

- Using the same VRF or network but using a different VRF name or network name in a different VXLAN fabric
- IPv6 underlay
- Security groups
- PVLAN
- TRM and TRMv6
- CloudSec VXLAN tunnel encryption
- Brownfield deployment where downstream VNI is already configured on a switch

Managing a brownfield VXLAN BGP EVPN fabric

This use case shows how to migrate an existing VXLAN BGP EVPN fabric to Cisco Nexus Dashboard. The transition involves migrating existing network configurations to Nexus Dashboard.

Typically, your fabric would be created and managed through manual CLI configuration or custom automation scripts. Now, you can start managing the fabric through Nexus Dashboard. After the migration, the fabric underlay and overlay networks will be managed by Nexus Dashboard.

Prerequisites

- Nexus Dashboard-supported NX-OS software versions. For details, refer to the Cisco Nexus Dashboard Release Notes.
- Underlay routing protocol is OSPF or IS-IS.
- The following fabric-wide loopback interface IDs must not overlap:
 - Routing loopback interface for IGP/BGP.
 - VTEP loopback ID
 - Underlay rendezvous point loopback ID if ASM is used for multicast replication.
- BGP configuration uses the 'router-id', which is the IP address of the routing loopback interface.
- If the iBGP peer template is configured, then it must be configured on the leaf switches and route reflectors. The template name that needs to be used between leaf and route reflector should be identical.
- The BGP route reflector and multicast rendezvous point (if applicable) functions are implemented on spine switches. Leaf switches do not support the functions.
- Familiarity with VXLAN BGP EVPN fabric concepts and functioning of the fabric from the Nexus Dashboard perspective.
- Fabric switch nodes are operationally stable and functional and all fabric links are up.
- vPC switches and the peer links are up before the migration. Ensure that no configuration updates are in progress or changes pending.
- Create an inventory list of the switches in the fabric with their IP addresses and credentials. Nexus Dashboard uses this information to connect to the switches.
- Shut down any other controller software you are using presently so that no further configuration changes are made to the VXLAN fabric. Alternatively, disconnect the network interfaces from the controller software (if any) so that no changes are allowed on the switches.
- The switch overlay configurations must have the mandatory configurations defined in the shipping Nexus Dashboard Universal Overlay profiles. Additional network or VRF overlay related configurations found on the switches are preserved in the freeform configuration associated with the network or VRF Nexus Dashboard entries.
- All the overlay network and VRF profile parameters such as VLAN name and route map name should be consistent across all devices in the fabric for the brownfield migration to be successful.

Guidelines and limitations

- Shared border fabric is not supported for brownfield migration.
- Brownfield import must be completed for the entire fabric by adding all the switches to the Nexus Dashboard fabric.
- The **cdp format device-id <system-name>** command to set the CDP device ID is not supported and will result in an error when adding switches during a brownfield import. The only supported format is **cdp format device-id <serial-number>** (the default format).

- On the **Create Fabric** window, the **Advanced > Overlay Mode** fabric setting decides how the overlays will be migrated. If the default config-profile is set, then the VRF and Network overlay configuration profiles will be deployed to switches as part of the migration process. In addition, there will be diffs to remove some of the redundant overlay CLI configurations. These are non network impacting.
- From the **Overlay Mode** drop-down list, if CLI is set, then VRF and Network overlay configurations stay on the switch as-is with no or little changes to address any consistency differences.
- The brownfield import in Nexus Dashboard supports the simplified NX-OS VXLAN EVPN configuration CLIs. For more information, see [Cisco N9000 Series NX-OS VXLAN Configuration Guide, Release 10.2\(x\)](#).
- The following features are unsupported.
 - Super Spine roles
 - ToR
 - eBGP underlay
 - Layer 3 port channel
- Take a backup of the switch configurations and save them before migration.
- No configuration changes (unless instructed to do so in this document) must be made to the switches until the migration is completed. Else, significant network issues can occur.
- Migration to Cisco Nexus Dashboard Fabric Controller is only supported for Cisco N9000 switches.
- The Border Spine and Border Gateway Spine roles are supported for the brownfield migration.
- First, note the guidelines for updating the settings. Then update each VXLAN fabric settings as explained below:
 - Some values (BGP AS Number, OSPF, etc) are considered as reference points to your existing fabric, and the values you enter must match the existing fabric values.
 - For some fields (such as IP address range, VXLAN ID range), the values that are auto-populated or entered in the settings are only used for future allocation. The existing fabric values are honored during migration.
 - Some fields relate to new functions that may not exist in your existing fabric (such as advertise-pip). Enable or disable it as per your need.
 - At a later point in time, after the fabric transition is complete, you can update settings if needed.

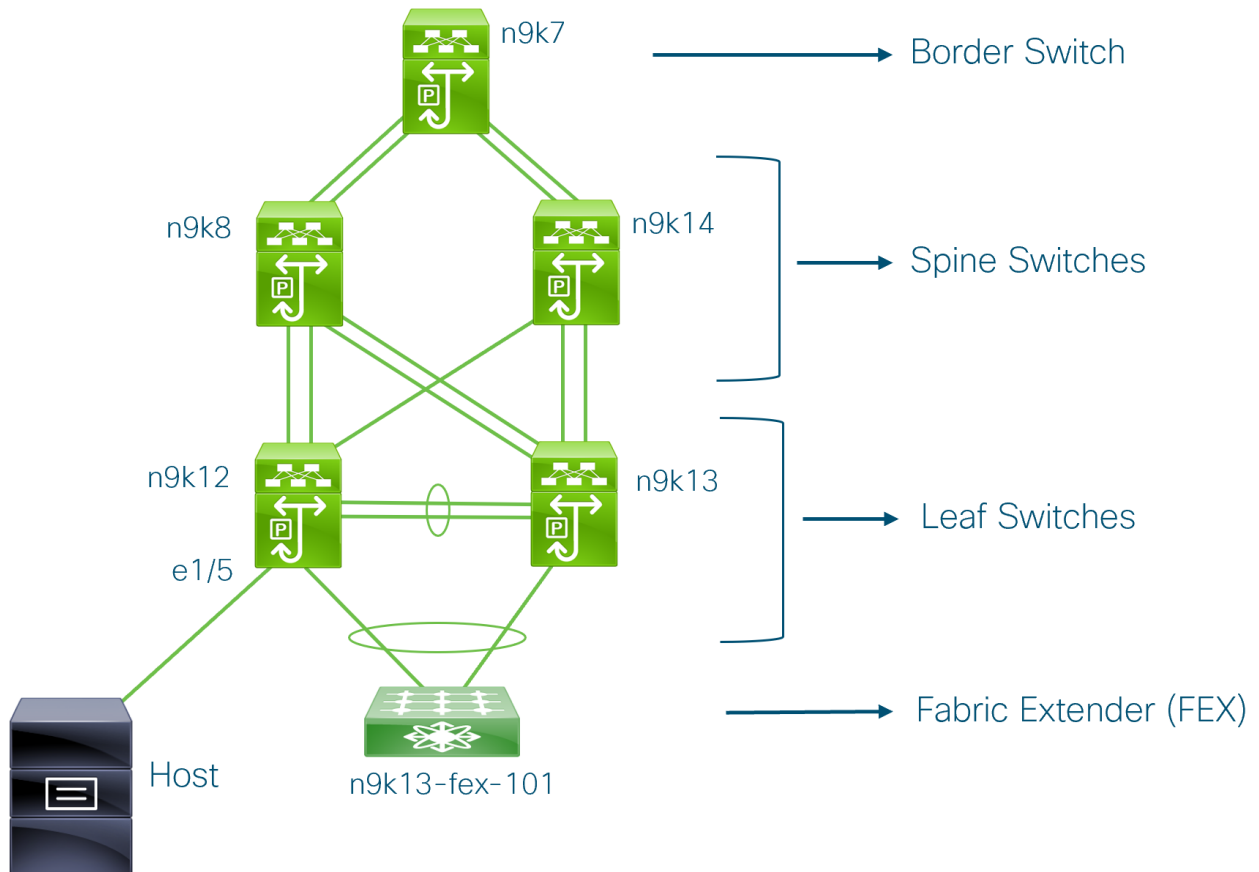
Fabric topology overview

This example use case uses the following hardware and software components:

- Five Cisco N9000 Series Switches
- One Fabric Extender or FEX
- One host

For information about the supported software images, see *Compatibility Matrix for Cisco Nexus Dashboard*.

Before we start the transition of the existing fabric, let us see its topology.



You can see that there is a border switch, two spine switches, two leaf switches, and a Fabric Extender or FEX.

A host is connected to the n9k12 leaf switch through the interface Ethernet 1/5.

Nexus Dashboard brownfield deployment tasks

The following tasks are involved in a Brownfield migration:

1. [Verify the existing VXLAN BGP EVPN fabric](#)
2. Create or edit an AI Data Center VXLAN EVPN fabric.
 - a. To create an AI Data Center VXLAN EVPN fabric, see [Creating Fabrics and Fabric Groups](#).
 - b. To edit an AI Data Center VXLAN EVPN fabric, see [Editing AI Data Center VXLAN fabric settings](#).

Verify the existing VXLAN BGP EVPN fabric

Let us check the network connectivity of the **n9k12** switch from the console terminal.

1. Verify the Network Virtual Interface or NVE in the fabric.

```
n9k12# show nve vni summary
```

Codes: CP - Control Plane DP - Data Plane

UC - Unconfigured

Total CP VNIs: 84 [Up: 84, Down: 0]

Total DP VNIs: 0 [Up: 0, Down: 0]

There are 84 VNIs in the control plane and they are up. Before the Brownfield migration, make sure that all the VNIs are up.

2. Check consistency and failures of vPC.

```
n9k12# show vpc
```

Legend:

(*) - local vPC is down, forwarding via vPC peer-link

```
vPC domain id          : 2
Peer status             : peer adjacency formed ok
vPC keep-alive status   : peer is alive
Configuration consistency status : success
Per-vlan consistency status : success
Type-2 consistency status : success
vPC role                : secondary
Number of vPCs configured : 40
Peer Gateway            : Enabled
Dual-active excluded VLANs : -
Graceful Consistency Check : Enabled
Auto-recovery status     : Enabled, timer is off.(timeout = 300s)
Delay-restore status     : Timer is off.(timeout = 60s)
Delay-restore SVI status : Timer is off.(timeout = 10s)
Operational Layer3 Peer-router : Disabled
.
.
.
```

3. Check the EVPN neighbors of the **n9k-12** switch.

```
n9k12# show bgp l2vpn evpn summary
```

```
BGP summary information for VRF default, address family L2VPN EVPN
BGP router identifier 192.168.0.4, local AS number 65000
BGP table version is 637, L2VPN EVPN config peers 2, capable peers 2
243 network entries and 318 paths using 57348 bytes of memory
BGP attribute entries [234/37440], BGP AS path entries [0/0]
BGP community entries [0/0], BGP clusterlist entries [2/8]
```

```
Neighbor    V  AS MsgRcvd MsgSent  TblVer  InQ OutQ Up/Down  State/PfxRcd
```

192.168.0.0	4	65000	250	91	637	0	0	01:26:59	75
192.168.0.1	4	65000	221	63	637	0	0	00:57:22	75

You can see that there are two neighbors corresponding to the spine switches.

Note that the ASN is 65000.

4. Verify the VRF information.

```
n9k12# show run vrf internet
```

```
!Command: show running-config vrf Internet
```

```
!Running configuration last done at: Fri Aug 9 01:38:02 2019
```

```
!Time: Fri Aug 9 02:48:03 2019
```

```
version 7.0(3)I7(6) Bios:version 07.59
```

```
interface Vlan347
  vrf member Internet
```

```
interface Vlan349
  vrf member Internet
```

```
interface Vlan3962
  vrf member Internet
```

```
interface Ethernet1/25
  vrf member Internet
```

```
interface Ethernet1/26
  vrf member Internet
vrf context Internet
  description Internet
  vni 16777210
  ip route 204.90.141.0/24 204.90.140.129 name LC-Networks
  rd auto
  address-family ipv4 unicast
    route-target both auto
    route-target both auto evpn
router ospf 300
  vrf Internet
    router-id 204.90.140.3
    redistribute direct route-map allow
    redistribute static route-map static-to-ospf
router bgp 65000
```

```
vrf Internet
  address-family ipv4 unicast
    advertise l2vpn evpn
```

The VRF **Internet** is configured on this switch.

The host connected to the **n9k-12** switch is part of the VRF **Internet**.

You can see the VLANs associated with this VRF.

Specifically, the host is part of **Vlan349**.

5. Verify the layer 3 interface information.

```
n9k12# show run interface vlan349
```

```
!Command: show running-config interface Vlan349
!Running configuration last done at: Fri Aug  9 01:38:02 2019
!Time: Fri Aug  9 02:49:27 2019
```

```
version 7.0(3)I7(6) Bios:version 07.59
```

```
interface Vlan349
  no shutdown
  vrf member Internet
  no ip redirects
  ip address 204.90.140.134/29
  no ipv6 redirects
  fabric forwarding mode anycast-gateway
```

Note that the IP address is **204.90.140.134**. This IP address is configured as the anycast gateway IP.

6. Verify the physical interface information. This switch is connected to the Host through the interface Ethernet 1/5.

```
n9k12# show run interface ethernet1/5
```

```
!Command: show running-config interface Ethernet1/5
!Running configuration last done at: Fri Aug  9 01:38:02 2019
!Time: Fri Aug  9 02:50:05 2019
```

```
version 7.0(3)I7(6) Bios:version 07.59
```

```
interface Ethernet1/5
  description to host
```

```
switchport mode trunk
switchport trunk native vlan 349
switchport trunk allowed vlan 349,800,815
spanning-tree bpduguard enable
mtu 9050
```

You can see that this interface is connected to the host and is configured with VLAN 349.

7. Verify the connectivity from the host to the anycast gateway IP address.

```
host# ping 204.90.140.134 count unlimited interval 1
PING 204.90.140.134 (204.90.140.134): 56 data bytes
64 bytes from 204.90.140.134: icmp_seq=0 ttl=254 time=1.078 ms
64 bytes from 204.90.140.134: icmp_seq=1 ttl=254 time=1.129 ms
64 bytes from 204.90.140.134: icmp_seq=2 ttl=254 time=1.151 ms
64 bytes from 204.90.140.134: icmp_seq=3 ttl=254 time=1.162 ms
64 bytes from 204.90.140.134: icmp_seq=4 ttl=254 time=1.84 ms
64 bytes from 204.90.140.134: icmp_seq=5 ttl=254 time=1.258 ms
64 bytes from 204.90.140.134: icmp_seq=6 ttl=254 time=1.273 ms
64 bytes from 204.90.140.134: icmp_seq=7 ttl=254 time=1.143 ms
```

We let the ping command run in the background while we transition the existing brownfield fabric into Nexus Dashboard Fabric Controller.

Add switches and transition VXLAN fabric management to Nexus Dashboard

Let us discover and add switches to the newly created fabric.

1. Navigate to **Manage > Fabrics**.
2. Click on the newly-created fabric to view the fabric **Overview** screen.
3. Choose **Inventory > Switches**.
4. From the **Actions** drop-down list, select **Add Switches**.

The **Add Switches** window appears.

Similarly, you can add switches on **Topology** window. On Topology window, choose a fabric, right-click on a fabric and click **Add Switches**.

5. On the **Add Switches - Fabric** screen, enter the **Seed Switch Details**.

Enter the IP address of the switch in the **Seed IP** field. Enter the username and password of the switches that you want to discover.

By default, the value in the **Max Hops** field is **2**. The switch with the specified IP address and the switches that are 2 hops from it will be populated after the discovery is complete.

Make sure to check the **Preserve Config** check box. This ensures that the current configuration of the switches will be retained.

6. Click **Discover Switches**.

The switch with the specified IP address and switches up to two hops away (depending on the setting of Max Hops) from it are populated in the Scan Details section.

7. Check the check box next to the switches that have to be imported into the fabric and click **Import into fabric**.

It is best practice to discover multiple switches at the same time in a single attempt. The switches must be cabled and connected to the Nexus Dashboard server and the switch status must be manageable.

If switches are imported in multiple attempts, then please ensure that all the switches are added to the fabric before proceeding with the Brownfield import process.

8. Click **Import into fabric**.

The switch discovery process is initiated. The **Progress** column displays progress for all the selected switches. It displays **done** for each switch after completion.



You should not close the screen and try to import switches again until all selected switches are imported or an error message comes up.

If an error message comes up, close the screen. The fabric topology screen comes up. The error messages are displayed at the top-right part of the screen. Resolve the errors and initiate the import process again by clicking **Add Switches** in the **Actions** panel.

9. After a successful import, the progress bar shows **Done** for all the switches. Click **Close**.

After closing the window, the fabric topology window comes up again. The switch is in Migration Mode now, and the Migration mode label is displayed on the switch icons.

At this point, you must not try to add Greenfield or *new* switches. Support is not available for adding new switches during the migration process. It might lead to undesirable consequences for your network. However, you can add a new switch after the migration process is complete.

10. After all the network elements are discovered, they are displayed in the **Topology** window in a connected topology. Each switch is assigned the **Leaf** role by default.

The switch discovery process might fail for a few switches, and the Discovery Error message is displayed. However, such switches are still displayed in the fabric topology. You should remove such switches from the fabric (Right-click the switch icon and click **Discovery > Remove** from fabric), and import them again.

You should not proceed to the next step until all switches in the existing fabric are discovered in Nexus Dashboard.

If you choose the Hierarchical layout for display (in the Actions panel), the topology automatically gets aligned as per role assignment, with the leaf switches at the bottom, the spine switches connected on top of them, and the border switches at the top.



The supported roles for switches with the Cisco NX-OS Release 7.0(3)I4(8b) and 7.0(4)I4(x) images are Border Leaf, Border Spine, Leaf, and Spine

11. Select the switch, click **Actions > Set Role**. On the Select Role screen, select **Border** and click **Select**.

Similarly, set the **Spine** role for the **n9k-14** and **n9k-8** spine switches.



You need to manually create a vPC pairing when the L3 keep alive is configured on the switches. Otherwise, the vPC configuration is automatically picked up from the switches.

vPC Pairing – The vPC pairing must be done for switches where the Layer 3 vPC peer-keep alive is used. The vPC configuration is automatically picked up from the switches when the vPC peer keep alive is established through the management option. This pairing reflects in the GUI only after the migration is complete.

- a. Right-click the switch icon and click vPC Pairing to set a vPC switch pair.

The Select vPC peer screen comes up. It lists potential vPC peer switches.

- b. Select the appropriate switch and click OK. The fabric topology comes up again. The vPC pair is formed.



Check if you have added all switches from the current fabric. If you have missed adding switches, add them now. Once you are certain that you have imported all existing switches, move to the next step, the Save and Deploy option.

12. From the Fabric Overview **Actions** drop-down list, choose **Recalculate and deploy**.

When you click **Recalculate and deploy**, Nexus Dashboard obtains switch configurations and populates the state of every switch from the current running configuration to the current expected configuration, which is the intended state maintained in Nexus Dashboard.

If there are configuration mismatches, **Pending Config** column shows the number of lines of difference. Click on the Pending Config column to view the **Pending Config** and **Side-by-Side Comparison** of the running configuration. Click **Deploy** to apply the configurations.

After the migration of underlay and overlay networks, the **Deploy Configuration** screen comes up.



- o The brownfield migration requires best practices to be followed on the existing fabric such as maintain consistency of the overlay configurations.
- o The Brownfield migration may take some time to complete since it involves collecting the running configuration from switches, build the Nexus Dashboard configuration intent based on these, consistency checks etc.
- o Any errors or inconsistencies that are found during the migration is reported in fabric errors. The switches continue to remain in the Migration mode. You should fix these errors and complete the migration again by clicking **Deploy** until no errors are reported.

13. After the configurations are generated, review them by clicking the links in the **Preview Config** column.

We strongly recommend that you preview the configuration before proceeding to deploy it on the

switches. Click the Preview Configuration column entry. The **Preview Config** screen comes up. It lists the pending configurations on the switch.

The Side-by-side Comparison tab displays the running configuration and expected configuration side-by-side.

The **Pending Config** tab displays the set of configurations that need to be deployed on a switch in order to go from the current running configuration to the current expected or intended configuration.

The **Pending Config** tab may show many configuration lines that will be deployed to the switches. Typically, on a successful brownfield import, these lines correspond to the configuration profiles pushed to the switches for an overlay network configuration. Note that the existing network and VRF-related overlay configurations are not removed from the switches.

The configuration profiles are Nexus Dashboard required constructs for managing the VXLAN configurations on the switches. During the Brownfield import process, they capture the same information as the original VXLAN configurations already present on the switches. In the following image, the configuration profile with **vlan 160** is applied.



Config Preview - Switch 80.80.80.62



Pending Config

Side-by-side Comparison

```
configure profile Auto_Net_VNI20160_VLAN160
vlan 160
  vn-segment 20160
  name 0160-BP2_RD_SGWS_Client_VLAN161_
interface Vlan160
  vrf member rd
  no ip redirects
  no ipv6 redirects
  ip address 10.9.160.1/24
  fabric forwarding mode anycast-gateway
  no shutdown
interface nve1
  member vni 20160
  ingress-replication protocol bgp
evpn
  vni 20160 12
  rd auto
  route-target import auto
  route-target export auto
configure terminal
apply profile Auto_Net_VNI20160_VLAN160
configure terminal
configure profile Auto_Net_VNI20180_VLAN180
vlan 180
```

As part of the import process, after the configuration profiles are applied, the original CLI based configuration references will be removed from the switches. These are the 'no' CLIs that will be seen towards the end of the diffs. The VXLAN configurations on the switches will be persisted in the configuration profiles. In the following image, you can see that the configurations will be removed, specifically, **no vlan 160**.

The removal of CLI based configuration is allowed if the **Overlay Mode** is set to **config-profile**, and not CLI.

Pending Config Side-by-side Comparison

```

no vlan 160
no vlan 159
no vlan 158
no vlan 157
no vlan 156
no vlan 155
no vlan 154
no vlan 126
no vlan 125
no vlan 124
no vlan 122
no vlan 1141
no vlan 10
no interface Vlan9
no interface Vlan899
no interface Vlan84
no interface Vlan820
no interface Vlan819
no interface Vlan818
no interface Vlan817
no interface Vlan816
no interface Vlan815
no interface Vlan814
no interface Vlan813

```

The **Side-by-side Comparison** tab displays the Running Config and Expected Config on the switch.

14. Close the **Config Preview Switch** window after reviewing the configurations.
15. Click **Deploy Config** to deploy the pending configuration onto the switches.

If the **Status** column displays **FAILED**, investigate the reason for failure to address the issue.

The progress bar shows **100%** for each switch. After correct provisioning and successful configuration compliance, close the screen.

In the fabric topology screen that comes up, all imported switch instances are displayed in green color, indicating successful configuration. Also, the **Migration Mode** label is not displayed on any switch icon.

Nexus Dashboard has successfully imported a VXLAN-EVPN fabric.

Post-transitioning of VXLAN fabric management to Nexus Dashboard – This completes the transitioning process of VXLAN fabric management to Nexus Dashboard. Now, you can add new switches and provision overlay networks for your fabric. For details, refer the respective section in the Fabrics topic in the configuration guide.

Configuration profiles support for brownfield migration

Cisco Nexus Dashboard supports the Brownfield import of fabrics with VXLAN overlay provisioned with configuration profiles. This import process recreates the overlay configuration intent based on the configuration profiles. The underlay migration is performed with the usual Brownfield migration.

This feature can be used to recover your existing fabric when a Nexus Dashboard backup is not available to be restored. In this case, you must install the latest Nexus Dashboard release, create a fabric, and then import the switches into the fabric.

Note that this feature is not recommended for the Nexus Dashboard upgrade. For more information, see *Cisco Nexus Dashboard Installation and Upgrade Guide*.

The following are the guidelines for the support of configuration profiles:

- The Brownfield migration of configuration profiles is supported for the **Data Center VXLAN EVPN** template.
- The configuration profiles on the switches must be a subset of the default overlay **Universal** profiles. If extra configuration lines are present that are not part of the **Universal** profiles, unwanted profile refreshes will be seen. In this case, after you recalculate and deploy configuration, review the diffs using the **Side-by-side Comparison** feature and deploy the changes.
- Brownfield migration with switches having a combination of VXLAN overlay configuration profiles and regular CLIs is not supported. If this condition is detected, an error is generated, and migration is aborted. All the overlays must be with either configuration profiles or regular CLIs only.

Manually add PIM-BIDIR configuration for leaf or spine post brownfield migration

After brownfield migration, if you add new spine or leaf switches, you should manually configure the PIM-BIDIR feature.

The following procedure shows how to manually configure the PIM-BIDIR feature for a new Leaf or Spine:

1. Check the **base_pim_bidir_11_1** policies that are created for an RP added through the brownfield migration. Check the RP IP and Multicast Group used in each `ip pim rp-address_RP_IP_group-list_MULTICAST_GROUP_bidir` command.
2. Add respective **base_pim_bidir_11_1** policies from the **View/Edit Policies** window for the new Leaf or Spine, push the config for each **base_pim_bidir_11_1** policy.

Migrate interconnected VXLAN fabrics with border gateway switches

When you migrate existing interconnected VXLAN fabrics with border gateway switches into Nexus Dashboard, make sure to note the following guidelines:

- Uncheck all **Auto** IFC creation related fabric settings. Review the settings and ensure they are unchecked as follows:
 - **Data Center VXLAN EVPN** fabric
Uncheck **Auto Deploy Both** check box under **Resources** tab.
 - Interconnected VXLAN fabrics
Uncheck **Multi-Site Underlay IFC Auto Deployment Flag** check box under **DCI** tab.
- Underlay Multisite peering: The eBGP peering and corresponding routed interfaces for underlay extensions between sites are captured in **switch_freeform** and **routed_interfaces**, and optionally in the **interface_freeform** configs. This configuration includes all the global configs for multisite. Loopbacks for EVPN multisite are also captured via the appropriate interface templates.
- Overlay Multisite peering: The eBGP peering is captured as part of **switch_freeform** as the only relevant config is under router bgp.

- Overlays containing Networks or VRFs: The corresponding intent is captured with the profiles on the Border Gateways with **extension_type = MULTISITE**.

1. Create all the required fabrics including the Data Center VXLAN EVPN and Multi-Site Interconnect Network fabrics with the required fabric settings. Disable the Auto VRF-Lite options as mentioned above. For more information, refer to *Creating VXLAN EVPN Fabric* and *External Fabric* sections.
2. Import all the switches into all the required fabrics and set roles accordingly.
3. Click **Recalculate and deploy** in each of the fabrics and ensure that the Brownfield Migration process reaches the 'Deployment' phase. Now, do not click **Deploy Configuration**.
4. Create the **VXLAN EVPN Multi-Site** fabric with the required fabric settings and disable the **Auto MultiSite IFC** options as shown in Guidelines. For more information, see *Creating a VXLAN EVPN Multi-Site Fabric*.
5. Move all the member fabrics into the VXLAN EVPN Multi-Site. Do not proceed further till this step is completed successfully. For more information, see *Moving the Member1 Fabric Under VXLAN EVPN Multi-Site-Parent-Fabric*.



The Overlay Networks and VRFs definitions in each of the Easy Fabrics must be symmetric for them to get added successfully to the VXLAN EVPN Multi-Site. Errors will be reported if any mismatches are found. These must be fixed by updating the overlay information in the fabric(s) and added to the VXLAN EVPN Multi-Site.

6. Create all the Multisite Underlay IFCs such that they match the IP address and settings of the deployed configuration.



Additional interface configurations must be added to the Source/Destination interface freeform fields in the **Advanced** section as needed.

For more information, see *Configuring Multi-Site Overlay IFCs*.

7. Create all the Multisite Overlay IFCs such that they match the IP address and settings of the deployed configuration. You will need to add the IFC links. For more information, see *Configuring Multi-Site Overlay IFCs*.
8. If there are VRF-Lite IFCs also, create them as well.



If the Brownfield Migration is for the case where Configuration Profiles already exist on the switches, the VRF-Lite IFCs will be created automatically in Step #3.

9. If Tenant Routed Multicast (TRM) is enabled in the VXLAN EVPN Multi-Site fabric, edit all the TRM related VRFs and Network entries in VXLAN EVPN Multi-Site and enable the TRM parameters.

This step needs to be performed if TRM is enabled in the fabric. If TRM is not enabled, you still need to edit each Network entry and save it.

10. Now click **Recalculate and deploy** in the VXLAN EVPN Multi-Site fabric, but, do not click **Deploy Configuration**.

11. Navigate to each member fabric, click **Recalculate and deploy**, and then click **Deploy Configuration**.

This completes the Brownfield Migration. You can now manage all the networks or VRFs for BGWs by using the regular Nexus Dashboard Overlay workflows.

When you migrate an existing VXLAN EVPN Multi-Site fabric with border gateway switches (BGW) that has a Layer-3 port-channel for Underlay IFCs, make sure to do the following steps:



Ensure that the child fabrics are added into VXLAN EVPN Multi-Site before migrating an VXLAN EVPN Multi-Site fabric.

1. Click on appropriate child fabric and navigate to **Fabrics Interfaces** to view the BGW. Choose an appropriate Layer-3 port channel to use for underlay IFC.
2. On **Policy** column, choose **int_port_channel_trunk_host_11_1** from drop-down list. Enter the associated port-channel interface members and then click **Save**.
3. Navigate to the **Tabular view** of the VXLAN EVPN Multi-Site fabric. Edit layer-3 port link, choose the multisite underlay IFC link template, enter source and destination IP addresses. These IP addresses are the same as existing configuration values on the switches
4. Do the steps from step 7 to 11 from above procedure.

Configuring a VXLANv6 fabric

You can create a fabric with an IPv6-only underlay. The IPv6 underlay is supported only for the **Data Center VXLAN EVPN** fabric type.

Nexus Dashboard provides support for configuring all VXLAN EVPN fabric types with an IPv6 underlay, support for PIMv6, and TRMv6. For more information, see [Configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6](#).

In the IPv6 underlay fabric, intra-fabric links, routing loopback, vPC peer link SVI, and NVE loopback interface for VTEPs are configured with IPv6 addresses. EVPN BGP neighbor peering is also established using IPv6 addressing.

Guidelines and limitations for an IPv6 underlay

- IPv6 underlay is supported for the Cisco N9000 Series switches with Cisco NX-OS Release 9.3(1) or higher.
- VXLANv6 is only supported for Cisco N9332C, Cisco N9364C, and Cisco Nexus modules that end with EX, GX, FX, FX2, FX3, or FXP.



VXLANv6 is defined as a VXLAN fabric with an IPv6 underlay.

- In VXLANv6, the platforms supported on a spine are all Cisco N9000 Series and Cisco Nexus 3000 Series platforms.
- The overlay routing protocol supported for the IPv6 fabric is BGP EVPN.
- vPC with the physical multichassis EtherChannel trunk (MCT) feature is supported for the IPv6 underlay network in Nexus Dashboard. The vPC peer keep-alive option can be loopback or

management with an IPv4 or an IPv6 address.

- Brownfield migration is supported for VXLANv6 fabrics. Note that L3 vPC keep-alive using an IPv6 address is not supported for a brownfield migration. This vPC configuration is deleted after the migration. However, L3 vPC keep-alive using an IPv4 address is supported.
- DHCPv6 is supported for the IPv6 underlay network.
- The following features are not supported for a VXLAN IPv6 underlay:
 - ISIS, OSPF, and BGP authentication
 - Dual-stack underlay
 - vPC fabric peering
 - DCI SR-MPLS or MPLS-LDP handoff
 - BFD
 - Super spine switch roles
 - NGOAM

Create a VXLAN EVPN fabric with IPv6 underlay

This procedure shows how to create a VXLAN EVPN fabric with an IPv6 underlay. Note that only the fields for creating a VXLAN fabric with an IPv6 underlay are documented. For information about the remaining fields, see [Creating Fabrics and Fabric Groups](#).

Nexus Dashboard added support for configuring a Data Center VXLAN EVPN fabric and a BGP VXLAN fabric with an IPv6 PIMv6 underlay. For more information, see [Configuring VXLANv6 with a PIMv6 underlay and TRMv6 for a Data Center VXLAN EVPN fabric](#).

1. Choose **Manage > Fabrics**.
2. From the **Actions** drop-down list, choose **Create Fabric**.

The **Create Fabric** page appears.

3. Enter the name of the fabric in the **Fabric Name** field.
4. In the **Pick Fabric** field, choose **Data Center VXLAN EVPN** from the **Select Type of Fabric** drop-down list.
5. Click **Select**.
6. The **General Parameters** tab is displayed by default. The fields for configuring an IPv6 underlay in the **General Parameters** tab are the following:

Field	Description
BGP ASN	Enter the BGP AS number for the fabric. You can enter either the 2 byte BGP ASN or 4 byte BGP ASN.
Enable IPv6 Underlay	Check the Enable IPv6 Underlay check box.

Field	Description
Enable IPv6 Link-Local Address	Check the Enable IPv6 Link-Local Address check box to use the link local addresses in the fabric between leaf-spine and spine-border interfaces. If you check this check box, the Underlay Subnet IPv6 Mask field is not editable. By default, the Enable IPv6 Link-Local Address field is enabled.
Fabric Interface Numbering	An IPv6 underlay supports p2p networks only. Therefore, the Fabric Interface Numbering drop-down list is disabled.
Underlay Subnet IPv6 Mask	Specify the subnet mask for the fabric interface for IPv6 addresses.
Underlay Routing Protocol	Specify the IGP used in the fabric, that is, OSPF or IS-IS for VXLANv6.
Enable IPv6 Underlay	Check the Enable IPv6 Underlay check box.
Enable IPv6 Link-Local Address	Check the Enable IPv6 Link-Local Address check box to use the link local addresses in the fabric between leaf-spine and spine-border interfaces. If you check this check box, the Underlay Subnet IPv6 Mask field is not editable. By default, the Enable IPv6 Link-Local Address field is enabled. IPv6 underlay supports p2p networks only. Therefore, the Fabric Interface Numbering drop-down list is disabled.
Underlay Subnet IPv6 Mask	Specify the subnet mask for the fabric interface for IPv6 addresses.
Underlay Routing Protocol	Specify the IGP used in the fabric, that is, OSPF or IS-IS for VXLANv6.

- Navigate to the **Replication** tab. The fields for configuring an IPv6 underlay for the Multicast replication mode are the following:

Field	Description
Replication Mode	The mode of replication that is used in the fabric for BUM (Broadcast, Unknown Unicast, Multicast) traffic. The choices are Ingress replication or Multicast replication. When you choose Ingress replication, the multicast-related fields are disabled. When you choose Multicast replication, the Ingress replication fields are disabled. You can change the fabric setting from one mode to the other, if no overlay profile exists for the fabric. Choose Multicast as the replication mode for Tenant Routed Multicast (TRM) for IPv4 or IPv6. For more information for the IPv4 use case, see Editing AI Data Center VXLAN fabric settings. For more information for the IPv6 use case, see Configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6.

Field	Description
IPv6 Multicast Group Subnet	Enter an IPv6 multicast address with a prefix range of 112 to 126.
Enable IPv4 Tenant Routed Multicast (TRM)	Check this check box to enable Tenant Routed Multicast (TRM) with IPv4 that allows overlay IPv4 multicast traffic to be supported over EVPN/MVPN in VXLAN EVPN fabrics.
Enable IPv6 Tenant Routed Multicast (TRM)	Check the check box to enable Tenant Routed Multicast (TRM) with IPv6 that allows overlay IPv6 multicast traffic to be supported over EVPN/MVPN in VXLAN EVPN fabrics.
Default MDT IPv4 Address for TRM VRFs	The multicast address for Tenant Routed Multicast traffic is populated. By default, this address is from the IP prefix specified in the Multicast Group Subnet field. When you update either field, ensure that the address is chosen from the IP prefix specified in Multicast Group Subnet. For more information, see Overview of Tenant Routed Multicast.
Default MDT IPv6 Address for TRM VRFs	The multicast address for Tenant Routed Multicast traffic is populated. By default, this address is from the IP prefix specified in the IPv6 Multicast Group Subnet field. When you update either field, ensure that the address is chosen from the IP prefix specified in IPv6 Multicast Group Subnet. For more information, see Overview of Tenant Routed Multicast.
MVPN VRI ID Range	Use this field for allocating a unique MVPN VRI ID per vPC. This field is needed for the following use cases: ▪ If you configure a VXLANv4 underlay with a Layer 3 VNI without VLAN mode, and you enable TRMv4 or TRMv6. ▪ If you configure a VXLANv6 underlay, and you enable TRMv4 or TRMv6.  The MVPN VRI ID cannot be the same as any site ID within a multi-site fabric. The VRI ID has to be unique within all sites within an MSD.
Enable MVPN VRI ID Re-allocation	Enable this check box to generate a one-time VRI ID reallocation. Nexus Dashboard automatically allocates a new MVPN ID within the MVPN VRI ID range above for each applicable switch. Since this is a one-time operation, after performing the operation, this field is turned off.  Changing the VRI ID is disruptive, so plan accordingly.

8. Navigate to the **VPC** tab.

Field	Description
vPC Peer Keep Alive option	Choose management or loopback . To use IP addresses assigned to the management port and the management VRF, choose management . To use IP addresses assigned to loopback interfaces and a non-management VRF, choose underlay routing loopback with an IPv6 address for PKA. Both the options are supported for an IPv6 underlay.

9. Navigate to the **Protocols** tab.

Field	Description
Underlay Anycast Loopback Id	Specify the underlay anycast loopback ID for an IPv6 underlay. You cannot configure an IPv6 address as secondary, an additional loopback interface is allocated on each vPC device. Its IPv6 address is used as the VIP.

10. Navigate to the **Resources** tab.

Field	Description
Manual Underlay IP Address Allocation	Check this check box to manually allocate underlay IP addresses. The dynamic underlay IP addresses fields are disabled.
Underlay Routing Loopback IPv6 Range	Specify the loopback IPv6 addresses for protocol peering.
Underlay VTEP Loopback IPv6 Range	Specify the loopback IPv6 addresses for VTEPs.
Underlay Subnet IPv6 Range	Specify the IPv6 address range that is used for assigning IP addresses for numbered and peer link SVIs. To edit this field, uncheck the Enable IPv6 Link-Local Address check box under the General Parameters tab.
BGP Router ID Range for IPv6 Underlay	Specify the address range to assign BGP Router IDs. The IPv4 addressing is used for routers with BGP and underlay routing protocols.

11. Navigate to the **Bootstrap** tab.

Field	Description
Enable Bootstrap	Check the Enable Bootstrap check box. If this check box is not chosen, none of the other fields on this tab are editable.
Enable Local DHCP Server	Check the check box to initiate automatic assignment of IP addresses assignment through the local DHCP server. The DHCP Scope Start Address and the DHCP Scope End Address fields are editable only after you check this check box.
DHCP Version	Choose DHCPv4 from the drop-down list.

12. Navigate to the **Advanced** tab.

Field	Description
Allow L3VNI w/o VLAN	Check this check box to allow Layer 3 VNI configuration without having to configure a VLAN.
Enable TCAM Allocation	Check this option to provide TCAM allocation to 768 or above if you enable TRMv6. TCAM commands are automatically generated for VXLAN and vPC fabric peering when enabled. For more information, see Configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6.

13. Click **Save** to complete the creation of the fabric.

What's next: See the section "Add switches to a fabric" in [Working with Inventory in Your Nexus Dashboard LAN or IPFM Fabrics](#).

Configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6

About creating VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6

Nexus Dashboard supports the following:

- Configuring VXLAN EVPN fabrics with Protocol Independent Multicast (PIM) IPv6 in the underlay (PIMv6)
- Configuring Tenant Routed Multicast (TRM) IPv6 for forwarding multicast traffic between senders and receivers within the same or different subnets or across Virtual Tunnel Endpoints (VTEPs)

This feature is supported when creating or editing the following fabric types:

- Data Center VXLAN EVPN fabric
- BGP (eBGP EVPN) fabric
- VXLAN EVPN Multi-Site fabric

The following sections provide information about configuring VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6:

- [Benefits of creating VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6](#)
- [Supported roles for configuring VXLAN fabrics with a PIMv6 underlay and TRMv6](#)
- [Supported platforms for VXLAN EVPN fabrics with a PIMv6 underlay](#)
- [Supported platforms for TRMv6](#)
- [Guidelines and limitations for VXLANv6 with a PIMv6 underlay and TRMv6](#)
- [Configuring VXLANv6 with a PIMv6 underlay and TRMv6 for a Data Center VXLAN EVPN fabric](#)
- [Configuring TRMv4 or TRMv6](#)

Benefits of creating VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6

- Supports multi-site fabrics
- Supports routing of IPv6 multicast traffic

- Supports an underlay type of IPv4 or IPv6 in a member fabric within a VXLAN fabric group
- Supports TRM IPv6
- Supports all BGW roles
- Supports IPv4 or IPv6 addresses
- Supports link local and BGP numbered

Supported roles for configuring VXLAN fabrics with a PIMv6 underlay and TRMv6

This table lists the supported roles for configuring VXLAN fabrics with a PIMv6 underlay and TRMv6.

Fabric Type	Roles
Data Center VXLAN EVPN Fabric	Leaf, Spine, Border, Border Gateway, Border Spine, Border Gateway Spine, ToR
BGP Fabric	Leaf, Spine, Border, Border Gateway, Border Spine, Border Gateway Spine, Super Spine, Border Super Spine, Border Gateway Super Spine

Supported platforms for VXLAN EVPN fabrics with a PIMv6 underlay

- Cisco NX-OS version 10.4.3 for BGWs
- Cisco NX-OS version 1.4.2 for anything other than a BGW
- Cisco N9300 FX, FX2, FX3, GX, GX2, and Cisco N9332D-H2R ToR switches are supported as the leaf VTEP
- Cisco N9500 with X9716D-GX line cards are supported only on the spine (EoR)
- Cisco N9500 with X9736C-FX line cards are supported only on the spine (EoR)



End of Row (EoR) requires configuring a non-default template using one of the following commands in global configuration mode: **system routing template-multicast-heavy** and **system routing template-multicast-ext-heavy**.

- Fabric peering (VMCT) is not supported with an IPv6 multicast underlay. VMCT is supported for an IPv6 ingress replication (IR) underlay beginning with Cisco NX-OS 10.3.2 on Cisco N9300 EX/FX/FX2/FX3/GX/GX2 switches and Cisco NX-OS 10.4.1 on a Cisco N9332D-H2R switch.

Supported platforms for TRMv6

- TRMv6
 - Cisco N9300 EX
 - Cisco N9300 FX, FX2, FX3, GX, and GX2 ToRs
- BGW
 - Cisco N9300 FX3, GX, and GX2 ToRs

Guidelines and limitations for VXLANv6 with a PIMv6 underlay and TRMv6

VXLANv6 does not support the following:

- No support for a vPC border gateway.

- No support for vPC fabric peering.
- No support for bidirectional forwarding for underlay multicast RP mode.
- No support for Bidirectional Forwarding Detection (BFD).
- Mp support for PIMv6 hello authentication.
- No support for any type of super spine role in a Data Center VXLAN EVPN fabric due to a restriction of an IPv6 underlay with IR.
- CloudSec
- No support for a Private VLAN (PVLAN).
- No support for a mix of IPv4 and IPv6 in OR3F.
- IPv6 multi-site loopback interface and Cisco Data Center Interconnect (DCI) interface address subnet pool is based on the underlay types.
- You cannot change a child fabric between IPv4 and IPv6 if the child fabric is part of VXLAN fabric group.
- Changing IPv4 to IPv6 addresses if the child fabric is part of a VXLAN fabric group.
- DCI uses ingress replication (IR).
- VXLANv6 with IR as an underlay is not supported for a VXLAN fabric group.

Configuring VXLANv6 with a PIMv6 underlay and TRMv6 for a Data Center VXLAN EVPN fabric

Follow these steps to configure VXLANv6 with a PIMv6 underlay and TRMv6 for a Data Center VXLAN EVPN fabric.

1. Create or edit a Data Center VXLAN EVPN fabric for an IPv6 underlay. For more information, see the section "Creating VXLAN EVPN fabrics with an IPv6 underlay" in [Editing Data Center VXLAN Fabric Settings](#).
2. Navigate to the **General Parameters** tab and check the **Enable IPv6 Underlay** and **Enable IPv6 Link-Local Address** options. For more information, see [General Parameters](#).
3. Navigate to the **Replication** tab and enable the following fields if you want to configure TRMv4 or TRMv6:
 - **Enable IPv4 Tenant Routed Multicast (TRM)**
 - **Enable IPv6 Tenant Routed Multicast (TRMv6)**
 - **MVPN VRI ID Range**

For more information, see [Replication](#).

4. Navigate to the **Advanced** tab and enable the **Allow L3 VNIw/o VXLAN** option.
5. Check the **Enable TCAM Allocation** option if you enable TRMv6.



A TCAM region needs to be allocated for TRMv6, which requires a reload of attached switches.

For more information, see [Advanced](#).

6. Click **Save** to complete the creation of the fabric.

Configuring TRMv4 or TRMv6

Follow these steps to configure TRMv4 or TRMv6.

1. Create or edit a VRF. For more information, see the "Create a VRF" section in [Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric](#).
2. Navigate to the **TRM** tab and enter the required fields depending on if you are enabling TRMv4 or TRMv6.



A TCAM region needs to be allocated for TRMv6, which requires a reload of the attached switches.

3. Click **Create** to create the VRF or click **Close** to discard the VRF.

Overview of VXLAN EVPN to SR-MPLS and MPLS LDP interconnection

Nexus Dashboard supports the following handoff features:

- VXLAN to SR-MPLS
- VXLAN to MPLS LDP

These features are provided on the border devices, that is, border leaf, border spine, and border super spine in the VXLAN fabric using the **Data Center VXLAN EVPN** template. Note that the devices should be running Cisco NX-OS Release 9.3(1) or later. These DCI handoff approaches are the one box DCI solution where no extra Provider Edge (PE) device is needed in the external fabric.



If the switch is running a Cisco NX-OS Release 7.0(3)I7(X), enabling the MPLS handoff feature causes the switch to remove the NVE related config-profile CLIs when the switch is reloaded.

In the Nexus Dashboard DCI MPLS handoff feature, the underlay routing protocol to connect a border device to an external fabric is ISIS or OSPF, and the overlay protocol is eBGP. The N-S traffic between the VXLAN fabric and external fabric running SR-MPLS or MPLS LDP is supported. Though, you can use Nexus Dashboard for connecting two Data Center VXLAN fabrics via SR-MPLS or MPLS LDP.

Supported platforms and configurations

The following table provides information about the supported platforms:

Feature	Supported Platforms
VXLAN to SR-MPLS	Cisco N9300-FX2/FX3/GX, N9K-X96136YC-R, and Cisco Nexus 3600 R-Series switches
VXLAN to MPLS LDP	N9K-X96136YC-R and Cisco Nexus 3600 R-series switches

The following features aren't supported as they aren't supported on a switch:

- Coexisting of MPLS LDP and SR-MPLS interconnections

- vPC

The VXLAN to SR-MPLS handoff feature comprises the following configurations:

- Base SR-MPLS feature configuration.
- Underlay configuration between the DCI handoff device and the device in the external fabric for the underlay connectivity. Nexus Dashboard supports ISIS or OSPF as the routing protocol for the underlay connectivity.
- Overlay configuration between a DCI handoff device and a core or edge router in the external fabric, or another border device in another fabric. The connectivity is established through eBGP.
- VRF profile

The VXLAN to MPLS LDP handoff feature comprises the following configurations:

- Base MPLS LDP feature configuration.
- Underlay configuration between the DCI handoff device and the device in the external fabric for the underlay connectivity. Nexus Dashboard supports ISIS or OSPF as the routing protocol for the underlay connectivity.
- Overlay configuration between a DCI handoff device and a core or edge router in the external fabric, or another border device in another fabric. The connectivity is established through eBGP.
- VRF profile

Inter-fabric connections for MPLS handoff

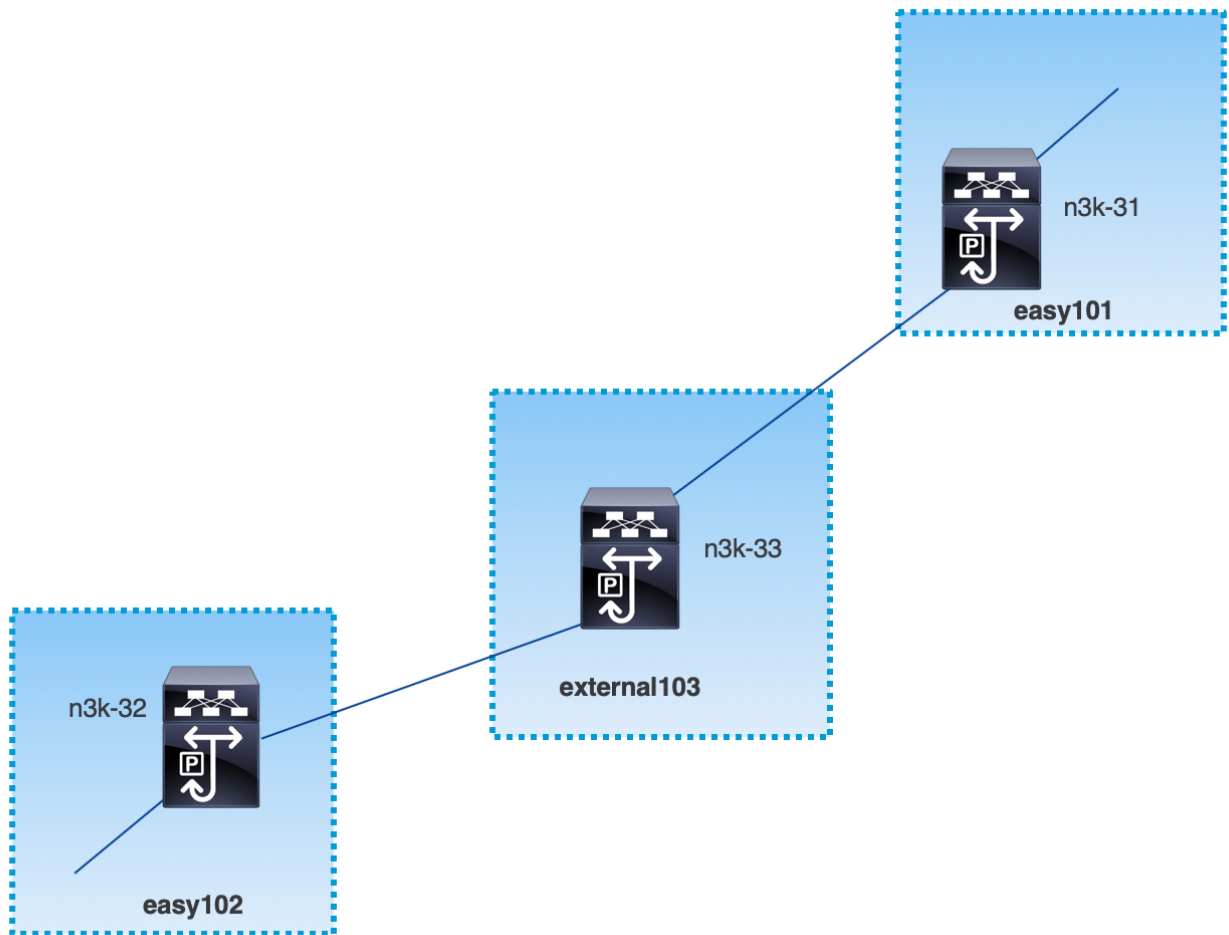
The following two inter-fabric connection links are available:

- **VXLAN_MPLS_UNDERLAY** for underlay configuration: This link corresponds to each physical link or Layer 3 port channel between the border and the external device (or a P router in MPLS or SR-MPLS). A border device can have multiple inter-fabric connection links as there could be multiple links connected to one or more external devices.
- **VXLAN_MPLS_OVERLAY** for eBGP overlay configuration: This link corresponds to the virtual link between a DCI handoff device and a core or edge router in the external fabric, or another border device in another fabric. This inter-fabric connection link can only be created on border devices which meet the image and platform requirement. A border device can have multiple of this type of IFC link as it could communicate to multiple core or edge routers.

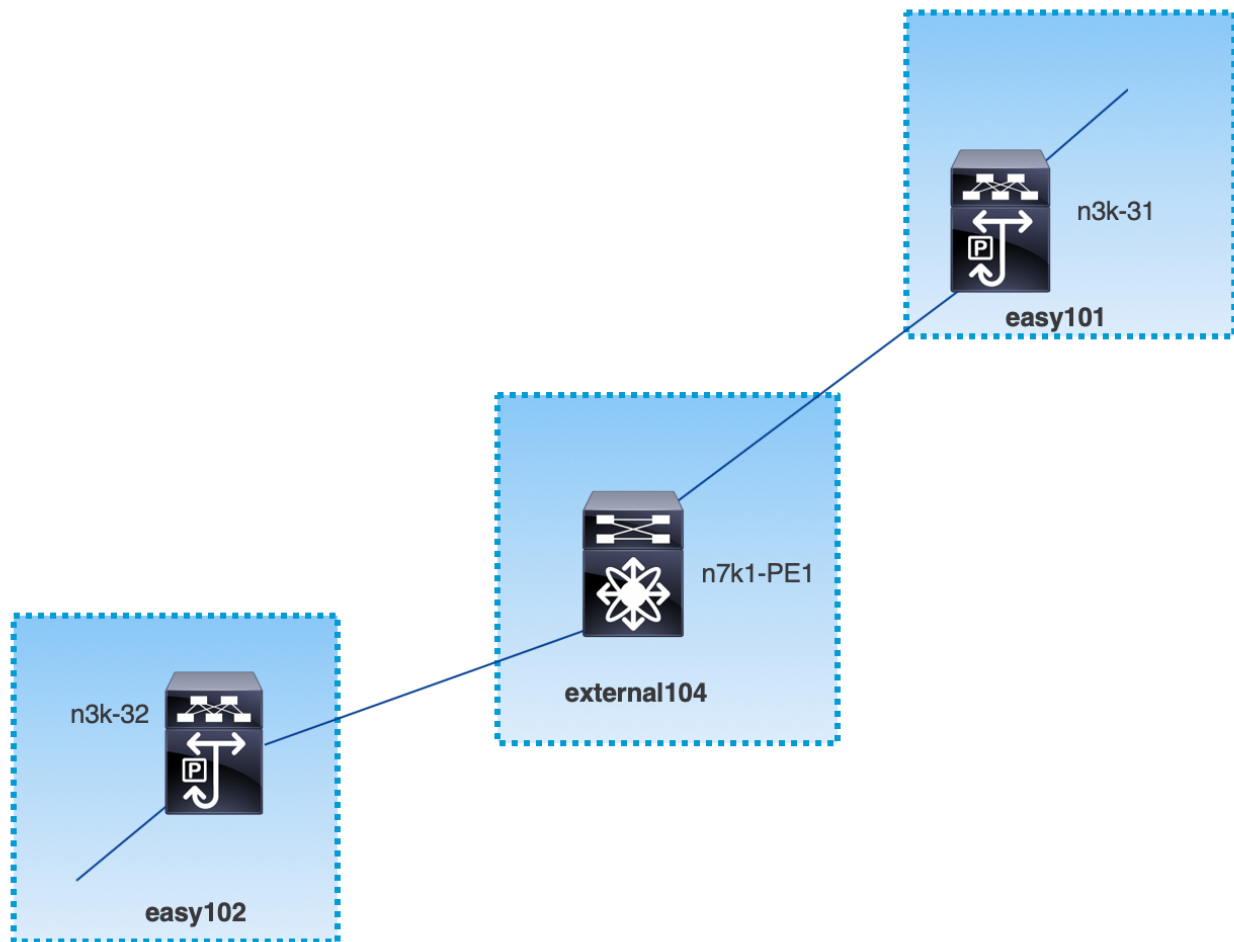
These inter-fabric connections can be manually created by using the Nexus Dashboard Web UI or REST API. Note that the automatic creation of these inter-fabric connections isn't supported.

VXLAN MPLS topology

MPLS-SR topology



MPLS-LDP topology



This topology shows only the border devices in the Data Center VXLAN EVPN and the core or edge router in the External Connectivity Network fabric.

- The fabrics that are using the **Data Center VXLAN EVPN** template are:
 - **easy101**
 - **easy102**
- The fabrics that are using the **External Connectivity Network** template are:
 - **external103**
 - **external104**
- The external fabric **external103** is running the MPLS SR protocol.
- The external fabric **external104** is running the MPLS LDP protocol.
- **n3k-31** and **n3k-32** are border devices performing VXLAN to MPLS handoff.
- **n7k-PE1** only supports MPLS LDP.
- **n3k-33** supports SR-MPLS.

Configuration tasks for VXLAN MPLS handoff

The following tasks are involved in configuring the MPLS handoff features:

1. Editing the fabric settings to enable MPLS handoff.

2. Creating an underlay inter-fabric connection link between the fabrics.

Specify whether you're using MPLS SR or LDP in the inter-fabric connection link settings.

3. Creating an overlay inter-fabric connection link between the fabrics.
4. Deploying a VRF for VXLAN to MPLS interconnection.

Editing fabric settings for MPLS handoff

This section shows how to edit the fabric settings for the VXLAN fabric and the external fabric to enable the MPLS handoff feature.

Edit VXLAN fabric settings

Follow these steps to edit VXLAN fabric settings.

1. Choose **Manage > Fabrics**.
2. Choose an appropriate VXLAN fabric.
3. Choose **Actions > Edit Fabric Settings** to edit the fabric settings.
4. Click **Fabric Management**.
5. Click **Advanced** and locate the following fields related to the MPLS handoff feature:
 - **Enable MPLS Handoff**: Check the box to enable the MPLS handoff feature.



When enabling the MPLS handoff feature for a brownfield import, most of the IFC configuration will be captured in **switch_freeform**.

- **Underlay MPLS Loopback Id**: Specify the underlay MPLS loopback ID.

The default value is 101. Valid range is 0 - 1023.

- **IS-IS NET Area Number for MPLS Handoff**: Specify the IS-IS NET area number for the MPLS handoff.
 - The value entered in this field is used only if the routing protocol on the DCI MPLS link is IS-IS
 - Enter the NET information in the form of **XX.<4-hex-digit Custom Area Number>.XXXX.XXXX.XXXX.00**
 - The default area number is 0001

6. Click **Resources** and locate the following field related to the MPLS handoff feature:
 - **Underlay MPLS Loopback IP Range**: Specify the underlay MPLS loopback IP address range.

For eBGP between border devices of VXLAN fabrics, the underlay routing loopback and underlay MPLS loopback IP range must be a unique range. VPNv4 peering will not come up if it overlaps with IP ranges of the other fabrics.

7. Click **Save** to configure the MPLS feature on each border device in the fabric.
8. In the **Fabrics** page, click the VXLAN fabric.

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

9. From **Actions** drop-down list, choose **Recalculate and deploy**.

Edit external fabric settings

Follow these steps to edit external fabric settings.

1. Choose **Manage > Fabrics**.
2. Choose an appropriate external fabric.
3. Choose **Actions > Edit Fabric Settings** to edit the fabric settings.
4. Click **Fabric Management**.
5. Click **General Parameters** and locate the following fields related to the MPLS handoff feature:
 - **Fabric Monitor Mode**: Uncheck the box in this field to disable the fabric monitor mode on the external fabric.
6. Click **Advanced** and locate the following fields related to the MPLS handoff feature:
 - **Enable MPLS Handoff**: Check the box to enable the MPLS handoff feature.
 - **Underlay MPLS Loopback Id**: Specify the underlay MPLS loopback ID.

The default value is 101. Valid range is 0 - 1023.

7. Click **Resources** and locate the following field related to the MPLS handoff feature:
 - **Underlay MPLS Loopback IP Range**: Specify the underlay MPLS SR or LDP loopback IP address range.

The IP range should be unique (it should not overlap with IP ranges of the other fabrics).

8. Click **Save** to configure the MPLS feature on each edge or core router in the fabric.
9. In the **Fabrics** page, click the external fabric.

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

10. Choose **Actions > Recalculate and deploy**.

Create an underlay inter-fabric connection

This procedure shows how to create an underlay inter-fabric connection link.

Follow these steps to create an underlay inter-fabric connection.

1. Choose **Manage > Fabrics**.
2. Choose a VXLAN fabric where you want to create an underlay inter-fabric connection to MPLS.
3. In the **Fabric Overview** page, choose **Connectivity > Links**.
4. Check the existing links that are already discovered for the fabric.

In this example, the link from **easy101** to **external103** is already discovered.

5. Select the existing discovered link and click on **Actions > Edit**.

If a link isn't discovered, click on **Actions > Create** and provide all the details for adding an inter-

fabric link.

6. In the **Link Management - Edit Link** page, provide all the required information.

- **Link Type:** Choose **Inter-Fabric**.
- **Link Sub-Type:** Choose **VXLAN_MPLS_Underlay** from the drop-down list.
- **Link Template:** Choose **ext_vxlan_mpls_underlay_setup** from the drop-down list.

In the **General Parameters** tab, provide all the details.

- **IP Address/Mask:** Specify the IP address with mask for the source interface.
- **Neighbor IP:** Specify the IP address of destination interface.
- **MPLS Fabric:** Specify whether the external fabric is running SR or LDP.



MPLS SR and LDP can't coexist on a single device.

- **Source SR Index:** Specify a unique SID index for the source border. This field is disabled if you choose **LDP** in the **MPLS Fabric** field.
- **Destination SR Index:** Specify a unique SID index for the destination border. This field is disabled if you choose LDP for the **MPLS Fabric** field.
- **SR Global Block Range:** Specify the SR global block range. You need to have the same global block range across the fabrics. The default range is from 16000 to 23999. This field is disabled if you choose LDP for the **MPLS Fabric** field.
- **DCI Routing Protocol:** Specify the routing protocol used on the DCI MPLS underlay link. You can choose either **is-is** or **ospf**.
- **OSPF Area ID:** Specify the OSPF area ID if you choose OSPF as the routing protocol.
- **DCI Routing Tag:** Specify the DCI routing tag used for the DCI routing protocol.

7. Click **Save**.

8. In the **Fabric Overview** page, choose **Actions > Recalculate and deploy**.

9. In the **Deploy Configuration** page, click **Deploy Config**.

10. Navigate to the destination fabric from the **Fabrics** page and perform a **Recalculate and deploy** on the destination fabric.

Create an overlay inter-fabric connection

This procedure shows how to create an overlay inter-fabric connection after the underlay inter-fabric connection is created. The overlay inter-fabric connection is the same for MPLS SR and LDP because the overlay connection uses eBGP.

Follow these steps to create an overlay inter-fabric connection.

1. Choose **Manage > Fabrics**.
2. Choose a VXLAN fabric where you want to create an overlay inter-fabric connection.
3. In the **Fabric Overview** page, choose **Connectivity > Links**.
4. Choose **Actions > Create**.

5. In the **Link Management - Create Link** page, provide all the details.

- **Link Type:** Choose **Inter-Fabric**.
- **Link-Sub Type:** Choose **VXLAN_MPLS_OVERLAY** from the drop-down list.
- **Link Template:** Choose **ext_vxlan_mpls_overlay_setup** from the drop-down list.
- **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 the MPLS loopback interface. The IP address of the loopback interface will be used for overlay eBGP peering.
- **Destination Device** and **Destination Interface:** Choose the destination device and a loopback interface that connects to the source device.

In the **General Parameters** tab, provide all the details.

- **BGP Local ASN:** In this field, the AS number of the source device is autopopulated.
- **BGP Neighbor IP:** Fill up this field with the IP address of the loopback interface at the destination device for eBGP peering.
- **BGP Neighbor ASN:** In this field, the AS number of the destination device is autopopulated.

6. Click **Save**.

7. In the **Fabric Overview** page, choose **Actions > Recalculate and deploy**.

8. In the **Deploy Configuration** page, click **Deploy Config**.

9. Navigate to the destination fabric from the **Fabrics** page and perform a **Recalculate and deploy** on the destination fabric.



If there is only one MPLS overlay IFC link on the switch, you can remove it only when there's no VRF attached to either end of the MPLS overlay link.

Deploy VRFs

This procedure shows how to deploy VRFs for VXLAN to MPLS interconnection.



When you use the 4 byte ASN and auto route target is configured, the route target that is automatically generated is 23456:VNI. If two different VRFs in two different fabrics have the same VNI value, the route-target of the two VRFs would be the same due to auto route target and the value 23456 is always constant. For two fabrics connected via VXLAN MPLS handoff, this could result in unintended route exchange. Therefore, for security reasons, if you want to disable auto route target, you can disable it by customizing the network template and network extension template.

Follow these steps to deploy VRFs.

1. Choose **Manage > Fabrics**.
2. Choose a VXLAN fabric.

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

3. Choose **Segmentation and security > VRFs**.
4. Click on **Actions > Create** to create a VRF.

For more information on creating VRFs, see the section "Working with VRFs" in [Working with Segmentation and Security for Your Nexus Dashboard VXLAN Fabric](#).

5. Choose the newly added VRF and click **Continue**.
6. In the **VRF Deployment** page, you can see the topology of the fabric. Select a border device to attach a VRF to the border device where the MPLS LDP IFC link is created.

In this example, **n3k-31** is the border device in the **easy101** fabric.

7. In the **VRF Extension Attachment** page, choose the VRF and click the **Freeform config** button under the CLI Freeform column.
8. Add the following freeform config manually to the VRF:

```
vrf context _$$VRF_NAME$$_  
  address-family ipv4 unicast  
    route-target import _$$REMOTE_PE_RT$$_  
  address-family ipv6 unicast  
    route-target import _$$REMOTE_PE_RT$$$
```

In the freeform config, *REMOTE_PE_RT* refers to the neighbor's BGP ASN and VNI number in the **ASN:VNI** format if the neighbor is a border device in a VXLAN fabric managed by Nexus Dashboard.

9. Click **Save Config**.
10. (Optional) Enter the Loopback Id and Loopback IPv4 Address and IPv6 address for the border device.
11. Click **Save**.
12. (Optional) Click the **Preview** icon in the **VRF Deployment** page to preview the configuration that will be deployed.
13. Click **Deploy**.

Perform the same task from Step 3 to Step 11 in the destination fabric if the neighbor is a border device in a VXLAN fabric managed by Nexus Dashboard.

Change the routing protocol and MPLS settings

This procedure shows how to change the routing protocol of a device from using IS-IS to OSPF, or from using MPLS SR to LDP for underlay IFC.



MPLS SR and LDP cannot co-exist on a device, and using both IS-IS and OSPF for MPLS handoff on the same device is not supported.

1. Remove all the MPLS underlay and overlay IFCs from the device that needs the change of DCI routing protocol or MPLS fabric.

2. Click **Recalculate and deploy** for each fabric that is involved in the removal of the IFCs.

This step deletes all global MPLS SR/LDP configurations and the MPLS loopback interface that was previously created.

3. Create a new IFC using the preferred DCI routing protocol and MPLS settings. For more information, see [Create an underlay inter-fabric connection](#).

Overview of Tenant Routed Multicast

Tenant Routed Multicast (TRM) enables multicast forwarding on the VXLAN fabric that uses a BGP-based EVPN control plane. TRM provides multi-tenancy aware multicast forwarding between senders and receivers within the same or different subnet local or across VTEPs.

With TRM enabled, multicast forwarding in the underlay is leveraged to replicate VXLAN encapsulated routed multicast traffic. A Default Multicast Distribution Tree (Default-MDT) is built per-VRF. This is an addition to the existing multicast groups for Layer-2 VNI Broadcast, Unknown Unicast, and Layer-2 multicast replication group. The individual multicast group addresses in the overlay are mapped to the respective underlay multicast address for replication and transport. The advantage of using a BGP-based approach allows the VXLAN BGP EVPN fabric with TRM to operate as fully distributed Overlay Rendezvous-Point (RP), with the RP presence on every edge-device (VTEP).

A multicast-enabled data center fabric is typically part of an overall multicast network. Multicast sources, receivers, and multicast rendezvous points might reside inside the data center but also might be inside the campus or externally reachable via the WAN. TRM allows a seamless integration with existing multicast networks. It can leverage multicast rendezvous points external to the fabric. Furthermore, TRM allows for tenant-aware external connectivity using Layer-3 physical interfaces or subinterfaces.

Guidelines and limitations

Refer to the following documents for switch-level guidelines and limitations for Tenant Routed Multicast:

- [Guidelines and Limitations for Tenant Routed Multicast](#)
- [Guidelines and Limitations for Layer 3 Tenant Routed Multicast](#)

Following are additional guidelines and limitations at the Nexus Dashboard level:

- When you perform a brownfield import on a fabric where the VRFs and the networks are deployed with the **Enable Tenant Routed Multicast** option enabled (without using a configuration profile), if the **Overlay Mode** option for the fabric is set to **cli**, the VRF-level **IPv4 TRM Enable** or the **IPv6 TRM Enable** option will not be enabled for VRFs imported into this fabric (in addition to the VRF-level **RP Address**, **RP Loopback ID**, **Underlay Mcast Address**, and **Overlay Mcast Groups** options, which will also not be enabled in this case).

In addition, if this VRF is deployed on a new leaf switch, IPv4 TRM or IPv6 TRM will not be enabled on that leaf switch for that VRF.

Overview of Tenant Routed Multicast with interconnected VXLAN fabrics

Tenant Routed Multicast (TRM) with interconnected VXLAN fabrics enables multicast forwarding across multiple VXLAN EVPN fabrics.

The following two use cases are supported:

- Use Case 1: TRM provides Layer 2 and Layer 3 multicast services across sites for sources and receivers across different sites.
- Use Case 2: Extending TRM functionality from a VXLAN fabric to source receivers external to the fabric.

Tenant Routed Multicast with interconnected VXLAN fabrics is an extension of the BGP-based TRM solution that enables multiple TRM sites with multiple VTEPs to connect to each other to provide multicast services across sites in most efficient possible way. Each TRM site operates independently and the border gateway on each site allows stitching across each site.

You can have multiple border gateways (BGWs) for each site. In a given site, the BGW peers with the route server or BGWs of other sites to exchange EVPN and MVPN routes. On the BGW, BGP imports routes into the local VRF/L3VNI/L2VNI and then advertises those imported routes into the fabric or WAN depending on where the routes were learnt from.

Operations for Tenant Routed Multicast with interconnected VXLAN fabrics

The operations for Tenant Routed Multicast (TRM) with interconnected VXLAN fabrics are as follows:

- Each fabric is represented by Anycast VTEP border gateways (BGWs). Designated forwarder (DF) election across BGWs ensures no packet duplication.
- Traffic between border gateways uses an ingress replication mechanism. Traffic is encapsulated with a VXLAN header followed by an IP header.
- Each fabric receives only one copy of the packet.
- Multicast source and receiver information across fabrics is propagated by BGP protocol on the border gateways configured with TRM.
- Border gateways on each fabric receives the multicast packet and re-encapsulates the packet before sending it to the local fabric.

For information about guidelines and limitations for TRM with interconnected VXLAN fabrics, see [Configuring Tenant Routed Multicast](#).

Configure Tenant Routed Multicast for a single site

This section assumes that a VXLAN EVPN fabric has already been provisioned through Nexus Dashboard.

Perform the following steps to enable Tenant Routed Multicast for a single site.

1. Enable Tenant Routed Multicast for a VXLAN fabric:
 - a. Navigate to **Manage > Fabrics**, then click the appropriate Data Center VXLAN EVPN fabric.

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

- b. Choose **Actions > Edit Fabric Settings**.
- c. Click **Fabric Management**.
- d. Click **Replication** and configure these fields:

Field	Description
Enable IPv4 Tenant Routed Multicast (TRM)	Check this check box to enable Tenant Routed Multicast (TRM) with IPv4 that allows overlay IPv4 multicast traffic to be supported over EVPN/MVPN in VXLAN EVPN fabrics.
Enable IPv6 Tenant Routed Multicast (TRM)	Check this check box to enable Tenant Routed Multicast (TRM) with IPv6 that allows overlay IPv6 multicast traffic to be supported over EVPN/MVPN in VXLAN EVPN fabrics.
Default MDT IPv4 Address for TRM VRFs	The multicast address for Tenant Routed Multicast traffic is populated. By default, this address is from the IP prefix specified in the Multicast Group Subnet field. When you update either field, ensure that the address is chosen from the IP prefix specified in Multicast Group Subnet .
Default MDT IPv6 Address for TRM VRFs	The multicast address for Tenant Routed Multicast traffic is populated. By default, this address is from the IP prefix specified in the IPv6 Multicast Group Subnet field. When you update either field, ensure that the address is chosen from the IP prefix specified in IPv6 Multicast Group Subnet .

- e. Click **Save** to save the fabric settings.

At this point, all the switches go into the pending state (the blue color).

- f. In the fabric's **Overview** page, choose **Actions > Recalculate and deploy** and then choose **Deploy Config** to enable the following:

- Enable feature ngMVPN: Enables the Next-Generation Multicast VPN (ngMVPN) control plane for BGP peering.
- Configure ip multicast multipath s-g-hash next-hop-based: Multipath hashing algorithm for the TRM-enabled VRFs.
- Configure ip igmp snooping vxlan: Enables IGMP Snooping for VXLAN VLANs.
- Configure ip multicast overlay-spt-only: Enables the MVPN Route-Type 5 on all MPVN-enabled Cisco N9000 switches.
- Configure and establish MVPN BGP AFI peering: This is necessary for the peering between BGP RR and the leaf nodes.

2. For a VXLAN EVPN fabric created using the eBGP overlay routing protocol, choose **Actions > Edit Fabric Settings** and navigate to **Fabric Management > EVPN** to enable the following fields, depending on if you are enabling IPv4 or IPv6:

- **Enable IPv4 Tenant Routed Multicast (TRM)** or **Enable IPv6 Tenant Routed Multicast (TRM)**
- **Default MDT Address for TRM VRFs** or **Default MDT IPv6 Address for TRM VRFs**

3. Enable Tenant Routed Multicast for the VRF:

- a. Navigate to **Segmentation and security > VRFs**.

- b. Choose the appropriate VRF and choose **Actions > Edit** to edit the selected VRF.
- c. Click **TRM** and edit these Tenant Routed Multicast settings:

Field	Description
IPv4 TRM Enable	Check the check box to enable IPv4 Tenant Routed Multicast. If you enable IPv4 TRMv4, and provide the RP address, you must enter the underlay multicast address in the Underlay Mcast Address field.
NO RP	Check the check box to disable RP fields. You must enable IPv4 Tenant Routed Multicast to edit this check box. If you enable No RP, then the Is RP External, RP Address, RP Loopback ID, and Overlay Mcast Groups fields are disabled.
Is RP External	Check this check box if the RP is external to the fabric. If this check box is not checked, RP is distributed in every VTEP.
RP Address	Specifies the IP address of the RP.
RP Loopback ID	Specifies the loopback ID of the RP, if Is RP External is not enabled.
Underlay Multicast Address	Specifies the multicast address associated with the VRF. The multicast address is used for transporting multicast traffic in the fabric underlay.  The multicast address in the Default MDT Address for TRM VRFs field on the fabric settings page is auto-populated in this field. You can override this field if a different multicast group address should be used for this VRF.
Overlay Mcast Groups	Specifies the multicast group subnet for the specified RP. The value is the group range in the ip pim rp-address command. If the field is empty, 224.0.0.0/24 is used as the default.
TRMv6 Enable	Check this check box to enable IPv6 Tenant Routed Multicast.
TRMv6 No RP	Check this check box to disable RP fields in TRMv6 as only PIM-SSM is used.
Is TRMv6 RP External	Check this check box if the RP is external to the fabric in TRMv6.
TRMv6 RP Address	Enter the IPv6 address for TRMv6 RP.
Overlay IPv6 Mcast Groups	Specifies the IPv6 multicast group subnet for the specified RP. The value is the group range in the ipv6 pim rp-address command. If the field is empty, ff00::/8 is used as the default.
Enable MVPN inter-as	Check this check box to use the inter-AS keyword for the Multicast VPN (MVPN) address family routes to cross the BGP autonomous system (AS) boundary. This option is applicable if you enabled the Tenant Routed Multicast option.

4. Click **Save** to save the settings.

The switches go into the pending state (the blue color). These settings enable the following:

- o Enable PIM on L3VNI SVI.
- o Route-Target Import and Export for MVPN AFI.
- o RP and other multicast configuration for the VRF.
- o Loopback interface using the above RP address and RP loopback id for the distributed RP.

5. Enable Tenant Routed Multicast for the network:

- a. Navigate to **Segmentation and security > Networks**.
- b. Choose the appropriate network and choose **Actions > Edit** to edit the selected network.
- c. Click **TRM** and edit these Tenant Routed Multicast settings.
- d. Check the **IPv4 TRM Enable** check box to enable Tenant Routed Multicast for IPv4 or check the **TRMv6 Enable** check box to enable Tenant Routed Multicast for IPv6.
- e. Click **Save** to save the settings.

The switches go into the pending state (the blue color). The Tenant Routed Multicast settings enable the following:

- Enable PIM on the L2VNI SVI.
- Create a PIM policy **none** to avoid PIM neighborship with PIM routers within a VLAN. The **none** keyword is a configured route map to deny any IPv4 or IPv6 addresses to avoid establishing a PIM neighborship policy using anycast IP.

Configure Tenant Routed Multicast with interconnected VXLAN fabrics

This section assumes that interconnected VXLAN fabrics have already been deployed through Nexus Dashboard and Tenant Routed Multicast (TRM) needs to be enabled.

Follow these steps to configure Tenant Routed Multicast (TRM) for interconnected VXLAN fabrics.

1. Navigate to **Manage > Fabrics**, then click the appropriate Data Center VXLAN EVPN fabric.

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

2. Choose **Segmentation and security > VRFs**.
3. Choose the appropriate VRF and choose **Actions > Edit** to edit the selected VRF.
4. Click **TRM**.
5. Edit the TRM settings as described in Step 3 of [Configure Tenant Routed Multicast for a single site](#).
6. Click **Save**.

The switches go into the pending state (the blue color). These settings enable the following:

- o Enable feature ngmvpn: Enables the Next-Generation Multicast VPN (ngMVPN) control plane for BGP peering.
- o Enables PIM on L3VNI SVI.
- o Configures the L3VNI multicast address.

- Route-target import and export for MVPN AFI.
- RP and other multicast configurations for the VRF.
- Loopback interface for the distributed RP.
- Enable the multi-site BUM ingress replication method for extending the Layer 2 VNI.

7. Establish MVPN AFI between the BGWs, as follows:

- a. Navigate to the fabric group:

Manage > Fabrics > Fabric groups

- b. Choose the fabric group to open the fabric group **Overview** page.

- c. Choose **Connectivity > Links**.

The **Links** subtab is chosen by default.

- d. Filter it by the policy - **Overlays**.

8. Select and edit each overlay peering to enable TRM by checking the **IPv4 Enable TRM** or the **TRMv6 Enable** check box.

9. Click **Save** to save the settings.

The switches go into the pending state, that is, the blue color. The TRM settings enable MVPN peering between the BGWs, or BGWs and the route server.

Understanding enhanced analytics for AI fabrics

Enhanced analytics is available for workloads in AI routed and VXLAN fabrics within Nexus Dashboard. This enhancement provides end-to-end visibility and actionable insights for AI infrastructures by integrating job completion and GPU statistics with network statistics, providing a detailed overview of network topologies along with GPUs.

In Nexus Dashboard 4.3.1, enhanced analytics for AI fabrics adds server network interface card (NIC) statistics to AI job monitoring. Before this release, AI job monitoring correlated Slurm job data, GPU telemetry, and switch-side network telemetry. With this release, you can correlate job data with server NIC telemetry, including error, discard, bandwidth, RDMA NACK, CNP, operational status, and link speed counters. You can use this correlation to determine whether job degradation is related to server-side NIC saturation or errors, switch-side network conditions, or both.

You can also use AI job monitoring with an existing AI fabric that is onboarded in Nexus Dashboard as an AI Observability Fabric. For onboarded AI fabrics, Nexus Dashboard can show AI jobs, server NIC statistics, correlated anomalies, and topology information when the required Slurm, GPU, server NIC, and topology telemetry is available.

For server NIC visibility, Nexus Dashboard collects server NIC counters every 30 seconds and retains the data for 15 days. The collected data includes CRC errors, transmit and receive discards, bandwidth Rx and Tx, NACK Rx and Tx, CNP Rx and Tx, operational status, and link speed. Server-interface optics counters are not available from the NVIDIA DOCA Telemetry Service (DTS) in this release.

These enhanced analytics for AI fabrics include:

- An AI resources page, which provides enhanced analytics data on these AI resources:
 - AI clusters
 - Scalable units
 - Servers
 - Live jobs
 - GPUs

For more information on AI resources, see [View information on AI resources](#).

- An AI job dashboard that correlates job performance with network health (such as optics, congestion, and errors), server NIC health (such as CRC errors, discards, bandwidth, NACK, CNP, operational status, and link speed), and GPU health (such as utilization, memory, temperature, and power). This enables proactive anomaly detection, efficient troubleshooting of job degradation, and optimization of network link utilization. For more information on AI jobs, see [View AI job information](#).

In order to view job details for AI fabrics, you must first integrate Slurm (Simple Linux Utility for Resource Management), which is an open-source job scheduler used to manage resources on high-performance computing clusters. For more information, see [Working with Integrations in Your Nexus Dashboard](#).

AI endpoint and topology discovery

Nexus Dashboard release 4.2.1 introduces enhanced discovery capabilities to provide unparalleled visibility into AI workloads. This feature extends traditional network device discovery in Nexus Dashboard to include detailed host-level information, offering a foundational understanding of your AI infrastructure, from network fabric to individual GPU servers.

Nexus Dashboard automatically identifies and collects the following rich data from GPU servers connected to your fabric.

- Host-level resources—Information about network interface cards (NICs) and GPU health.
- Intelligent correlation—Multiple network interfaces (NICs/IP addresses) from a single physical server are intelligently correlated and presented as a unified host entity within Nexus Dashboard. This simplified management provides a holistic view of each server, even with multi-NIC configurations.
- Logical grouping—Nexus Dashboard organizes discovered hosts and their connected leaf switches into logical constructs called scalable units and AI Clusters, which are essential for structured management and analysis.

The core of this enhanced discovery relies on LLDP. When enabled for AI fabric types, Nexus Dashboard actively queries connected devices to gather LLDP details. This protocol provides key host attributes such as system name and description, local and remote port IDs, MAC and IP addresses, and management addresses.

For AI fabric types, Nexus Dashboard enables LLDP endpoint discovery on the switches by default. You must also enable LLDP on the hosts. Nexus Dashboard then automatically filters switches and routers based on endpoint descriptions, such as vendor names in the system description.

Nexus Dashboard introduces specific logical groupings, to effectively manage and visualize AI environments.

- Scalable unit—A logical grouping of leaf switches and servers. Nexus Dashboard calculates the operational status of a scalable unit based on the health of its constituent leaf switches and servers. Individual link statuses (e.g., between a leaf and a server/GPU) do not directly participate in this high-level scalable unit status calculation.
- AI cluster—An AI cluster is essentially a fabric group that provides a dedicated view and management space specifically for your AI infrastructure. You can associate a fabric with an AI cluster if you see the **Associate with AI cluster** option in that fabric, as described in [Associate a fabric with an AI cluster](#).

For more information, see [View information on AI resources](#).

Prerequisites

- The enhanced discovery is specifically available for AI fabric types (AI Routed and VXLAN EVPN fabrics) and the LLDP endpoint discovery is enabled by default for these fabric types. You must also enable LLDP on the hosts as well.
- To organize resources into scalable units, you must assign switches with the leaf role to a scalable unit through user-driven actions within AI fabrics. See [Associate a switch with a scalable unit](#) for more information.
- To view job details for AI fabrics, you must integrate Simple Linux Utility for Resource Management (Slurm) with Nexus Dashboard. For more information, see [Working with Integrations in Your Nexus Dashboard](#).
- To collect GPU statistics, DCGM (Nvidia Data Center GPU Manager) Exporter should be running on the host. See [DCGM Exporter](#) for more information.
- To collect server NIC statistics, NVIDIA DOCA Telemetry Service (DTS) and Fluent Bit must be configured on the GPU servers to send server NIC metrics to Nexus Dashboard.
- GPU servers must be reachable from the Nexus Dashboard data network.
- The hostname reported by the Data Center GPU Manager (DCGM) Exporter, Data Center Infrastructure on a Chip Architecture (DOCA) Telemetry Service, Slurm, and LLDP advertisements must match the server hostname.
- You must also run the script to collect GPU-to-NIC mapping information for AI fabric analytics. See [GPU/NIC Topology Fleet Collector](#) for more information.

Guidelines and limitations: Analytics for AI fabrics

These are the guidelines and limitations for the analytics for AI fabrics feature:

- Only **Config-Only** backup and restore operations are supported with this feature. Full backup and restore operations are not supported.
- This functionality is only available for co-hosted clusters.
- To view server NIC statistics, configure the GPU servers to expose NVIDIA DOCA Telemetry Service (DTS) metrics and use Fluent Bit to send the server NIC metrics to Nexus Dashboard.
- Server NIC counters are collected every 30 seconds and retained for 15 days.

- Server-interface optics counters are not available from DTS in this release.
- For external AI fabrics, AI job monitoring depends on successful fabric onboarding, Slurm integration, GPU and NIC topology mapping, and telemetry availability. If switch telemetry is not available, switch-side interface metrics are not displayed.
- Run the GPU/NIC Topology Fleet Collector scripts:
 - During the initial setup,
 - If you change the hostname of the server at any point after you have set up your AI fabrics, or
 - If you have any topology change, interface rename or NIC replacement.

Upload the generated zip file into the **Update version** box in the **Job monitoring telemetry** area under **Edit <fabric> Settings > AI settings**. See [GPU/NIC Topology Fleet Collector](#) and [AI settings](#) for more information.

- Nexus Dashboard pulls fresh jobs data:
 - Manually, whenever you click **Resync**
 - Automatically, every 5 minutes

Jobs started and completed within this window will not get correlated with GPU details.

- When you perform a backup and restore operation, after you have completed the restore process, complete jobs data is read from Slurm and is shown in job monitoring. Historical jobs will not get correlated with GPU details.
- The AI job **Topology** view does not retain the topology from the time the job ran. Instead, it displays the job topology using the currently available topology data and resource associations. If an AI cluster, fabric association, scalable unit, switch, server, NIC, or GPU changes after the job runs, the topology for that job might be incomplete.

Guidelines and limitations: GPU Anomalies

- All GPU anomalies are enabled by default on startup.
 - GPU anomaly levels: By default, only **Critical** (set at 90%) is enabled.
 - GPU memory/utilization/power: The critical value is set at 90%.
 - GPU temperature: The critical value is set at 90 degrees Celsius.
- You cannot set or modify the GPU power anomaly thresholds from the global alert thresholds page. Only these default settings are in effect:
 - Warning: 0
 - Major: 0
 - Critical: 90
- Internal thresholds are also set for the GPU temperature, which is based on the limits published by the DCGM Exporter. If those thresholds are breached, then an anomaly will be raised irrespective of the user's setting on the global alerts page.
- Warnings are always suppressed. There is no effect if you change the settings for the Warning levels.
- In the **GPU memory anomaly** area, there is a known issue where the **What's wrong** message

provides mismatched unit information.

Add Slurm as an integration

In order to view job details for AI fabrics, you must first integrate Slurm (Simple Linux Utility for Resource Management), which is an open-source job scheduler used to manage resources on high-performance computing clusters. For more information, see [Working with Integrations in Your Nexus Dashboard](#).

DCGM Exporter

To collect GPU statistics, DCGM (Nvidia Data Center GPU Manager) Exporter should be running on the host.

DCGM Exporter is a tool based on the Go APIs to NVIDIA DCGM that allows users to gather GPU metrics and understand workload behavior or monitor GPUs in clusters. DCGM Exporter exposes GPU metrics at an HTTP endpoint (/metrics) for monitoring solutions such as Prometheus.

To use the DCGM Exporter, see [DCGM Exporter](#).



- You may want to run the DCGM exporter in secure mode. For that, you must upload a CA certificate and a client certificate/key in the **Certificate Management** page and attach those certificates to the **AIComputeVisibility** feature. See [Collect data in secured mode](#) for more information.
- You should validate that the DCGM Exporter is properly configured. You can make a REST API call to the DCGM Exporter and validate that the hostname is correct and that the counters in the [Working with the dcp-metrics-included.csv file](#) section are present in the response.

Working with the dcp-metrics-included.csv file

The **dcp-metrics-included.csv** file has a list of counters to collect. You will need this file at the start of DCGM Exporter process.

The **dcp-metrics-included.csv** file must have these entries.

```
# Temperature
```

```
DCGM_FI_DEV_GPU_TEMP,    gauge, GPU temperature (in C).
```

```
# Utilization (the sample period varies depending on the product)
```

```
DCGM_FI_DEV_GPU_UTIL,    gauge, GPU utilization (in %).
```

```
# Power
```

```
DCGM_FI_DEV_POWER_USAGE, gauge, Power draw (in W).
```

```
#Slowdown temperature for the device.
```

DCGM_FI_DEV_SLOWDOWN_TEMP, gauge, Slowdown temperature for the device.

#Shutdown temperature for the device.

DCGM_FI_DEV_SHUTDOWN_TEMP, gauge, Shutdown temperature for the device.

#Current Power limit for the device.

DCGM_FI_DEV_POWER_MGMT_LIMIT, gauge, Current power limit for the device.

#Used Frame Buffer in MB.

DCGM_FI_DEV_FB_USED, gauge, Framebuffer memory used (in MiB).

#Free Frame Buffer in MB.

DCGM_FI_DEV_FB_FREE, gauge, Framebuffer memory free (in MiB).

#Total Frame Buffer of the GPU in MB.

DCGM_FI_DEV_FB_TOTAL, gauge, Framebuffer memory total (in MiB).

#Percentage used of Frame Buffer: 'Used/(Total - Reserved)'.

DCGM_FI_DEV_FB_USED_PERCENT, gauge, Frame buffer used (in percent).

DCGM_FI_DEV_COUNT, counter, Number of Devices on the node.

DCGM_FI_DEV_NAME, label, Name of device

DCGM_FI_DEV_BRAND, label, Device Brand

DCGM_FI_DEV_SERIAL, label, Device Serial Number

DCGM_FI_DEV_UUID, label, UUID

#Effective power limit that the driver enforces after taking into account all limiters.

DCGM_FI_DEV_ENFORCED_POWER_LIMIT, gauge, Enforced power limit for the device.

Configure DOCA Telemetry Service and Fluent Bit

Use the NVIDIA Data Center Infrastructure on a Chip Architecture (DOCA) Telemetry Service (DTS) and Fluent Bit on each GPU server to forward server network interface card (NIC) metrics to the

Nexus Dashboard OpenTelemetry collector. DTS exposes server NIC metrics on a local Prometheus endpoint. Fluent Bit scrapes the DTS endpoint every 30 seconds and forwards the metrics to the CollectorPersistent endpoint in Nexus Dashboard.

Prerequisites

- Ensure that Docker is installed on the DGX host.
- Ensure that the DGX host has network connectivity to the Nexus Dashboard cluster.
- Obtain the CollectorPersistent pod IP address and OTLP NIC receiver port from the Nexus Dashboard cluster.
- For secure communication with the OpenTelemetry (OTel) collector using mutual TLS (mTLS), ensure that the CA certificate, server certificate, and server key are available.
- Ensure that Docker is installed on each GPU server that sends server NIC metrics.
- Ensure that the GPU server has network reachability to the Nexus Dashboard data network.
- Retrieve the CollectorPersistent external data IP address from Admin > System Settings > External Data IPs in Nexus Dashboard.
- Use OpenTelemetry Protocol (OTLP) port 4319 for HTTP, or OTLP port 4320 for HTTPS with mutual TLS (mTLS).
- For HTTPS with mTLS, copy the CA certificate, Fluent Bit client certificate, and key to the GPU server.
- Ensure that the hostname reported by Data Center GPU Manager (DCGM) Exporter, DTS, Slurm, and Link Layer Discovery Protocol (LLDP) advertisements matches the server hostname.

Initial configuration

1. Pull and initialize the DOCA Telemetry Service (DTS) container on the DGX host:

```
export DTS_IMAGE=nvcr.io/nvidia/doca/doca_telemetry:1.21.5-doca3.0.0-host
docker run -v "/opt/mellanox/doca/services/telemetry/config:/config" \
  --rm --name doca-telemetry-init -it $DTS_IMAGE \
  /bin/bash -c "DTS_CONFIG_DIR=host /usr/bin/telemetry-init.sh"
```

The system creates the default DTS configuration files in </opt/mellanox/doca/services/telemetry/config/>.

2. Edit the /opt/mellanox/doca/services/telemetry/config/dts_config.ini file on the DGX host. Locate the **PROVIDERS** section and update it to enable the **sysfs** and **ethtool** providers while disabling unnecessary sub-components. Replace the existing provider lines with the following::

```
# >>> BEGIN – changed section >>>

enable-provider=sysfs
# the sysfs.rate component calculates rate of sysfs counters and is usually unnecessary
disable-provider=sysfs.rate
# the sysfs.ib_mr_cache component is usually unnecessary
```

```
disable-provider=sysfs.ib_mr_cache
#disable-provider=sysfs.eth
disable-provider=sysfs.hwmon

enable-provider=ethtool

# <<< END – changed section <<<
```

Ensure the Prometheus endpoint is enabled. Verify that the prometheus line is not commented out. c. Save the file.

```
prometheus=http://0.0.0.0:9100
```

Start or restart DOCA Telemetry Service (DTS) container to apply the updated configuration.



For the complete `dts_config.ini` file with these changes applied in Appendix A.

- Before you configure Fluent Bit, retrieve the following information from the Nexus Dashboard (ND) cluster:

Parameter	How to obtain	Example
collectorpersistent1 pod IP	This is the IP under admin → system settings → external data ips , (check the one assigned to collectorpersistent1 pod) on the ND UI.	10.10.10.5
OTLP NIC receiver port	HTTP: 4319 , HTTPS: 4320 (configured on the OTel collector)	Use 4319 for HTTP, 4320 for HTTPS
mTLS certificates	Generate a fresh bundle using Appendix B. Before generating, update the collector VIP in the script so the OTEL server certificate SAN matches the CollectorPersistent external IP.	Copy to /root/mtls_certs/ on the DGX host using the directory layout shown below

Certificate requirements

Option	Description
HTTP	No certificates are required on the DGX host.
HTTPS with mTLS	Copy the generated Fluent-side CA and client certificate files to the DGX host using the following layout:

```

mkdir -p /root/mtls_certs/ca
mkdir -p /root/mtls_certs/fluentbit_receiver

# Copy these files from the generated bundle:
# /aiinfra-cert/mtls_certs/ca/aiinfra_ca.crt
# /aiinfra-cert/mtls_certs/dcgm_exporter_server/dcgm_exporter_server.crt
# /aiinfra-cert/mtls_certs/dcgm_exporter_server/dcgm_exporter_server.key

cp /aiinfra-cert/mtls_certs/ca/aiinfra_ca.crt /root/mtls_certs/ca/aiinfra_ca.crt
cp /aiinfra-cert/mtls_certs/dcgm_exporter_server/dcgm_exporter_server.crt
/root/mtls_certs/fluentbit_receiver/dcgm_exporter_server.crt
cp /aiinfra-cert/mtls_certs/dcgm_exporter_server/dcgm_exporter_server.key
/root/mtls_certs/fluentbit_receiver/dcgm_exporter_server.key

```



For HTTPS with mTLS on port 4320, enable TLS verification in the Fluent Bit output and specify the CA certificate, client certificate, and client key files.

4. Create the Fluent Bit configuration. Follow these steps to create the Fluent Bit configuration file ([/root/fluent-bit-nic.yaml](#)) on the DGX host based on your transport requirements.
 - a. Non-TLS (Port 4319): If the collector exposes the NIC receiver on port 4319 without Transport Layer Security (TLS), use the following configuration:

```

pipeline:
  inputs:
    - name: prometheus_scrape
      tag: nic_rdma_metrics
      host: 127.0.0.1
      port: 9100
      scrape_interval: 30s
      metrics_path: /metrics
      buffer_max_size: 32M
      processors:
        metrics:
          - name: metrics_selector
            metric_name:
              /^(rx_crc_errors_phy|tx_discards_phy|rx_discards_phy|rx_bytes_phy|rx_packets_phy|rx_bytes|tx_bytes_phy|tx_packets_phy|tx_bytes|hw_rnr_nak_retry_err|hw_implied_nak_seq_err|hw_req_transport_retries_exceeded|hw_local_ack_timeout_err|hw_out_of_buffer|hw_out_of_sequence|port_rcv_constraint_errors|port_rcv_errors|local_link_integrity_errors|port_rcv_remote_physical_errors|symbol_error|port_rcv_switch_relay_errors|hw_rp_cnp_handled|hw_np_cnp_sent|hw_port_state|current_link_speed|max_link_speed|current_link_width|max_link_width|rx_global_pause_duration|tx_global_pause_duration|rx_global_pause_transition|tx_global_pause_transition|rx_prio1_pause_duration|tx_prio1_pause_duration|rx_prio2_pause_duration|tx_prio2_pause

```

```

_duration|rx_prio3_pause_duration|tx_prio3_pause_duration|rx_prio4_pause_duratio
n|tx_prio4_pause_duration|rx_prio5_pause_duration|tx_prio5_pause_duration|rx_pri
o6_pause_duration|tx_prio6_pause_duration|rx_prio7_pause_duration|tx_prio7_pau
se_duration)$/
    action: include
outputs:
  - name: opentelemetry
    match: nic_rdma_metrics
    host: <COLLECTOR_PERSISTENT_POD_IP>
    port: 4319
    metrics_uri: /v1/metrics

  - name: stdout
    match: '*'
    format: json_lines

```

- b. TLS with client authentication (Port 4320): If the collector exposes the NIC receiver on port 4320 with TLS client authentication enabled, use the following configuration:

```

pipeline:
  inputs:
    - name: prometheus_scrape
      tag: nic_rdma_metrics
      host: 127.0.0.1
      port: 9100
      scrape_interval: 30s
      metrics_path: /metrics
      buffer_max_size: 32M
      processors:
        metrics:
          - name: metrics_selector
            metric_name:
/^(rx_crc_errors_phy|tx_discards_phy|rx_discards_phy|rx_bytes_phy|rx_packets_ph
y|rx_bytes|tx_bytes_phy|tx_packets_phy|tx_bytes|hw_rnr_nak_retry_err|hw_implied
_nak_seq_err|hw_req_transport_retries_exceeded|hw_local_ack_timeout_err|hw_ou
t_of_buffer|hw_out_of_sequence|port_rcv_constraint_errors|port_rcv_errors|local_li
nk_integrity_errors|port_rcv_remote_physical_errors|symbol_error|port_rcv_switch_r
elay_errors|hw_rp_cnp_handled|hw_np_cnp_sent|hw_port_state|current_link_speed
|max_link_speed|current_link_width|max_link_width|rx_global_pause_duration|tx_gl
obal_pause_duration|rx_global_pause_transition|tx_global_pause_transition|rx_prio1
_pause_duration|tx_prio1_pause_duration|rx_prio2_pause_duration|tx_prio2_pause
_duration|rx_prio3_pause_duration|tx_prio3_pause_duration|rx_prio4_pause_duratio
n|tx_prio4_pause_duration|rx_prio5_pause_duration|tx_prio5_pause_duration|rx_pri
o6_pause_duration|tx_prio6_pause_duration|rx_prio7_pause_duration|tx_prio7_pau
se_duration)$/

```

```

    action: include
  outputs:
    - name: opentelemetry
      match: nic_rdma_metrics
      host: <COLLECTOR_PERSISTENT_POD_IP>
      port: 4320
      metrics_uri: /v1/metrics
      tls: on
      tls.verify: on
      tls.ca_file: /root/mtls_certs/ca/aiinfra_ca.crt
      tls.crt_file: /root/mtls_certs/fluentbit_receiver/dcgm_exporter_server.crt
      tls.key_file: /root/mtls_certs/fluentbit_receiver/dcgm_exporter_server.key

    - name: stdout
      match: '*'
      format: json_lines

```

- c. Create the `/root/fluent-bit-nic.yaml` file on the DGX host and paste the content corresponding to your selected transport method and save the file.
- d. Replace the placeholders in the template and select the configuration that matches your transport method:

Placeholder	Value
<COLLECTOR_PERSISTENT_POD_IP>	IP obtained in Step 3

5. Start the Fluent Bit container on the DGX host using the command that matches your chosen transport.
 - a. Update the `<path>` placeholder in the following commands to match the location of your certs folder and `fluent-bit-nic.yaml` file.
 - b. Start the Fluent Bit container using the command that matches your transport method:
 - i. Run Fluent-bit for HTTP.

```

docker run -d --name fluent-bit-nic \
  --network=host \
  -v <path>/fluent-bit-nic.yaml:/fluent-bit/etc/fluent-bit.yaml:ro \
  fluent/fluent-bit:4.2.3 \
  /fluent-bit/bin/fluent-bit -c /fluent-bit/etc/fluent-bit.yaml

```

- ii. Run Fluent-bit for HTTPS with mTLS.

```

docker run -d --name fluent-bit-nic \
  --network=host \
  -v <path>/fluent-bit-nic.yaml:/fluent-bit/etc/fluent-bit.yaml:ro \

```

```
-v /root/mtls_certs:/root/mtls_certs:ro \
fluent/fluent-bit:4.2.3 \
/fluent-bit/bin/fluent-bit -c /fluent-bit/etc/fluent-bit.yaml
```



For HTTPS with mTLS, also mount the certificate directory.

- iii. `--network=host`: Allows Fluent Bit to scrape `localhost:9100` (the DOCA Telemetry Service (DTS) Prometheus endpoint) and reach the collector pod IP address directly.
- iv. HTTPS with mTLS: You must mount both the Fluent Bit configuration file and the `/root/mtls_certs` directory as read-only.

Verification

Check the Fluent Bit logs to confirm that metrics are being scraped and forwarded correctly.

1. View the container logs: `docker logs -f fluent-bit-nic`
2. Verify the following indicators of a successful setup:
 - a. Successful startup: The logs show no errors regarding the configuration or TLS.
 - b. Metrics processing: Standard output (`stdout`) displays JSON-formatted NIC metrics.
 - c. Connection stability: There are no repeated connection-refused or TLS handshake errors on the OpenTelemetry output.
3. Verify the DTS Prometheus endpoint metrics: `curl -s http://localhost:9100/metrics | head -50`

The DTS endpoint must return server NIC metrics from the sysfs and eth tool providers, and the Fluent Bit logs must not show repeated connection or TLS errors.



If Nexus Dashboard does not receive server NIC metrics through Fluent Bit for more than 15 minutes, the system raises an anomaly. If this anomaly occurs, verify that the DOCA Telemetry Service (DTS) and Fluent Bit are running, confirm that the DTS metrics endpoint is returning data, and check the Fluent Bit logs for connection or TLS errors.

Appendix A: Full `dts_config.ini`

```
[clx]
##### GENERAL CONFIGURATION
#####
sync-time-limit=10000
counter-buffer-size=65536
event-buffer-size=65536
verbose=6
update=1000
standalone=true
runtime-configuration-folder=/config/clx_last_runtime_conf
idle-time-limit=0
```

```
##### PROVIDERS
#####
enable-explicitly=true

# --- CHANGED SECTION START ---
enable-provider=sysfs
disable-provider=sysfs.rate
disable-provider=sysfs.ib_mr_cache
#disable-provider=sysfs.eth
disable-provider=sysfs.hwmon

enable-provider=ethtool
# --- CHANGED SECTION END ---

#enable-provider=vnic
#enable-provider=ifconfig
#enable-provider=diagnostic_data_low_freq
#diagnostic-data-yml-file=/config/diagnostic_data_configs/all-single-port.yml
#enable-provider=amber
#amber_devices=DEV1,DEV2,DEV3
#amber_update_interval_sec=30
#enable-provider=ppcc_eth
#ppcc_eth_devices=mt4129_pciconf0,mt4129_pciconf0.1
#ppcc_algo_slot=0
#ppcc_algo_param_index=0
#ppcc_local_port=1
#ppcc_pnat=0
#ppcc_lp_msb=0
#enable-provider=nvidia_smi
#nvidia_smi_max_processes=20
#nvsmi_with_numeric_fields=1
#enable-provider=dcgm
dcgm_events_enable_common_fields=1
dcgm_events_enable_prof_fields=0
dcgm_events_enable_ecc_fields=0
dcgm_events_enable_nvlink_fields=0
dcgm_events_enable_nvswitch_fields=0
dcgm_events_enable_vgpu_fields=0
dcgm_events_profile_period=1000
dcgm_events_monitor_period=5000
dcgm_events_static_period=120000

##### DATA OUTPUTS
#####
```

```
#output=/data
schema-path=/data/schema
data-template={{year}}/{{month}}/{{day}}/{{source}}/{{tag}}{{id}}.bin
```

```
##### Prometheus
#####
prometheus=http://0.0.0.0:9100
prometheus-ignore-tags=FI_metrics
```

```
##### Fluent Bit
#####
fluentbit-config-dir=/config/fluent_bit_configs
```

```
##### IPC transport
#####
ipc-sockets-dir=/tmp/ipc_sockets
```

```
##### NetFlow
#####
#netflow-collector-addr=0
#netflow-collector-port=0
```

```
##### Open Telemetry
#####
#open-telemetry-receiver=http://64.100.43.47:4318/v1/metrics
```

```
##### Remote Write
#####
#remote-write-receiver=http://0.0.0.0:9090/api/v1/write
```

```
##### Real Time Analysis Exporter
#####
#lua-script-dir=/config/rta/lua_scripts
```

```
##### HTTP APIs
#####
#enable-http-api=true
```

```
##### Labels
#####
#level-labels-file=/config/level_labels.ini
```

```
##### Ad Hoc File
#####
```

```
ad-hoc-runner-file=/config/clx_ad_hoc_runner_config.ini
```

Appendix B: Generate new mTLS certificates

Use this script to generate a new CA, a Fluent Bit client certificate, and an OpenTelemetry (OTEL) server certificate for the HTTPS with mTLS flow.



Before running the script, update **OTEL_COLLECTOR_VIP_IP** to match the current **CollectorPersistent** external IP address. The SAN of the generated OTEL certificate must match the IP address to which Fluent Bit connects.

The script generates the certificate bundle locally in **/aiinfra-cert/mtls_certs**. Copy the generated CA and Fluent Bit client files to the DPU host paths specified in Step 3. You must also install the matching OTEL certificate and CA on the collector side.

```
#!/bin/bash
set -euo pipefail

# Base names for the certificate components
CA_BASE_NAME="aiinfra_ca"
DCGM_EXPORTER_BASE_NAME="dcgm_exporter_server"
OTEL_COLLECTOR_BASE_NAME="otelcollector_client"

# Validity in days
CA_VALIDITY_DAYS=3650
SERVER_VALIDITY_DAYS=365
CLIENT_VALIDITY_DAYS=365

# Common Name (CN) for the Root CA
CA_CN="Allinfra Root CA"

# DCGM Exporter Server Certificate Configuration
DCGM_EXPORTER_SERVICE_NAME="dcgm-exporter-service"
K8S_NAMESPACE="default"
K8S_CLUSTER_DOMAIN="svc.cluster.local"
DCGM_EXPORTER_CN="${DCGM_EXPORTER_SERVICE_NAME}.${K8S_NAMESPACE}.${K8S_CLUSTER_DOMAIN}"
DCGM_EXPORTER_SAN_HOSTS="DNS:${DCGM_EXPORTER_CN},DNS:${DCGM_EXPORTER_SERVICE_NAME},DNS:*.${DCGM_EXPORTER_SERVICE_NAME}.${K8S_NAMESPACE}.${K8S_CLUSTER_DOMAIN},DNS:*.cisco.com,IP:127.0.0.1,IP:64.100.43.47,IP:64.100.43.48,IP:64.100.43.49"

# OTEL Collector Dual-Use Certificate Configuration
OTEL_COLLECTOR_CN="otelcollector.cisco-nir.svc"
OTEL_COLLECTOR_VIP_IP="10.127.121.103"
OTEL_COLLECTOR_SAN_HOSTS="IP:${OTEL_COLLECTOR_VIP_IP},DNS:${OTEL_COLLECTOR_CN},DNS:*.cisco.com,IP:127.0.0.1,IP:64.100.43.47,IP:64.100.43.48,IP:64.100.43.49"
```

```
OR_CN}"
```

```
# Encryption algorithm for encrypted private keys
```

```
KEY_ENCRYPTION_ALGO=" aes256"
```

```
# Passphrases
```

```
CA_PASSPHRASE=" aiinfra"
```

```
OTEL_PASSPHRASE=" aiinfra"
```

```
# Output Directory
```

```
OUTPUT_DIR=" /aiinfra-cert"
```

```
rm -rf "$OUTPUT_DIR"
```

```
mkdir -p "$OUTPUT_DIR/mtls_certs/ca"
```

```
"$OUTPUT_DIR/mtls_certs/dcgm_exporter_server"
```

```
"$OUTPUT_DIR/mtls_certs/otelcollector_client"
```

```
CDIR="$OUTPUT_DIR/mtls_certs"
```

```
echo " Certificates will be generated in: $CDIR"
```

```
openssl genpkey -algorithm RSA -out "$CDIR/ca/${CA_BASE_NAME}.key.tmp"
```

```
openssl pkcs8 -topk8 -v2 "${KEY_ENCRYPTION_ALGO}" -in
```

```
"$CDIR/ca/${CA_BASE_NAME}.key.tmp" -out "$CDIR/ca/${CA_BASE_NAME}.key"
```

```
-passout pass:" $CA_PASSPHRASE"
```

```
rm "$CDIR/ca/${CA_BASE_NAME}.key.tmp"
```

```
openssl req -x509 -new -nodes -key "$CDIR/ca/${CA_BASE_NAME}.key" -sha256 -days
```

```
"$CA_VALIDITY_DAYS" -out "$CDIR/ca/${CA_BASE_NAME}.crt" -subj "/CN=${CA_CN}"
```

```
-passin pass:"$CA_PASSPHRASE"
```

```
openssl genpkey -algorithm RSA -out
```

```
"$CDIR/dcgm_exporter_server/${DCGM_EXPORTER_BASE_NAME}.key"
```

```
cat > "$CDIR/dcgm_exporter_server/server_ext.cnf" <<EOF2
```

```
[req]
```

```
distinguished_name = req_distinguished_name
```

```
req_extensions = v3_req
```

```
prompt = no
```

```
[req_distinguished_name]
```

```
C = US
```

```
ST = California
```

```
L = San Jose
```

```
O = MyOrg
```

```
OU = DCGM Exporter
```

```
CN = ${DCGM_EXPORTER_CN}
```

```
[v3_req]
```

```

subjectAltName = ${DCGM_EXPORTER_SAN_HOSTS}
extendedKeyUsage = clientAuth,serverAuth
EOF2
openssl req -new -key
"$CDIR/dcgm_exporter_server/${DCGM_EXPORTER_BASE_NAME}.key" -out
"$CDIR/dcgm_exporter_server/${DCGM_EXPORTER_BASE_NAME}.csr" -config
"$CDIR/dcgm_exporter_server/server_ext.cnf"
cat > "$CDIR/dcgm_exporter_server/v3_server.cnf" <<EOF2
authorityKeyIdentifier=keyid,issuer
basicConstraints=CA:FALSE
keyUsage=digitalSignature,keyEncipherment
extendedKeyUsage=clientAuth,serverAuth
subjectAltName=${DCGM_EXPORTER_SAN_HOSTS}
EOF2
openssl x509 -req -in
"$CDIR/dcgm_exporter_server/${DCGM_EXPORTER_BASE_NAME}.csr" -CA
"$CDIR/ca/${CA_BASE_NAME}.crt" -CAkey "$CDIR/ca/${CA_BASE_NAME}.key"
-CAcreateserial -out
"$CDIR/dcgm_exporter_server/${DCGM_EXPORTER_BASE_NAME}.crt" -days
"$SERVER_VALIDITY_DAYS" -sha256 -extfile
"$CDIR/dcgm_exporter_server/v3_server.cnf" -passin pass:"$CA_PASSPHRASE"

openssl genpkey -algorithm RSA -out
"$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.key.tmp"
openssl pkcs8 -topk8 -v2 "${KEY_ENCRYPTION_ALGO}" -in
"$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.key.tmp" -out
"$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.key" -passout
pass:"$OTEL_PASSPHRASE"
rm "$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.key.tmp"
cat > "$CDIR/otelcollector_client/client_ext.cnf" <<EOF2
[req]
distinguished_name = req_distinguished_name
req_extensions = v3_req
prompt = no

[req_distinguished_name]
C = US
ST = California
L = San Jose
O = Allinfra
OU = OTel Collector
CN = ${OTEL_COLLECTOR_CN}

[v3_req]
subjectAltName = ${OTEL_COLLECTOR_SAN_HOSTS}

```

```

extendedKeyUsage = clientAuth,serverAuth
EOF2
openssl req -new -key
"$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.key" -out
"$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.csr" -config
"$CDIR/otelcollector_client/client_ext.cnf" -passin pass:"$OTEL_PASSPHRASE"
cat > "$CDIR/otelcollector_client/v3_client.cnf" <<EOF2
authorityKeyIdentifier=keyid,issuer
basicConstraints=CA:FALSE
keyUsage=digitalSignature,keyEncipherment
extendedKeyUsage=clientAuth,serverAuth
subjectAltName=${OTEL_COLLECTOR_SAN_HOSTS}
EOF2
openssl x509 -req -in "$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.csr"
-CA "$CDIR/ca/${CA_BASE_NAME}.crt" -CAkey "$CDIR/ca/${CA_BASE_NAME}.key"
-CAcreateserial -out "$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.crt"
-days "$CLIENT_VALIDITY_DAYS" -sha256 -extfile
"$CDIR/otelcollector_client/v3_client.cnf" -passin pass:"$CA_PASSPHRASE"

rm "$CDIR/dcgm_exporter_server/server_ext.cnf"
"$CDIR/dcgm_exporter_server/v3_server.cnf" "$CDIR/otelcollector_client/client_ext.cnf"
"$CDIR/otelcollector_client/v3_client.cnf"

printf '\n--- VERIFY CHAIN ---\n'
openssl verify -CAfile "$CDIR/ca/${CA_BASE_NAME}.crt"
"$CDIR/dcgm_exporter_server/${DCGM_EXPORTER_BASE_NAME}.crt"
openssl verify -CAfile "$CDIR/ca/${CA_BASE_NAME}.crt"
"$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.crt"
printf '\n--- OTEL CERT DETAILS ---\n'
openssl x509 -in "$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.crt" -text
-noout | sed -n '/X509v3 Subject Alternative Name/,+3p;/X509v3 Extended Key
Usage/,+2p'
printf '\n--- KEY ENCRYPTION CHECK ---\n'
if grep -q 'BEGIN ENCRYPTED PRIVATE KEY'
"$CDIR/otelcollector_client/${OTEL_COLLECTOR_BASE_NAME}.key"; then echo
'otel_key_encrypted=yes'; else echo 'otel_key_encrypted=no'; fi
if grep -q 'BEGIN ENCRYPTED PRIVATE KEY'
"$CDIR/dcgm_exporter_server/${DCGM_EXPORTER_BASE_NAME}.key"; then echo
'dcgm_key_encrypted=yes'; else echo 'dcgm_key_encrypted=no'; fi
printf '\n--- FILES ---\n'
find "$CDIR" -maxdepth 2 -type f | sort

```

GPU/NIC Topology Fleet Collector

The GPU/NIC Topology Fleet Collector is a toolkit to collect one time information from AI servers for

AI infrastructure visibility. On Nexus Dashboard, this information is used to derive physical connections between GPUs and NICs on a given server, which helps in correlating GPU and network issues.

The GPU/NIC Topology Fleet Collector collects:

- GPU information (model, UUID, PCI addresses)
- Network Interface Cards (NICs) with specifications
- GPU↔NIC topology distances.
- Network interface details (IPs, MACs, drivers)
- DCGM monitoring deployment status

Prerequisites

- Before you run the GPU/NIC Topology Fleet Collector script, verify that the DCGM Exporter, NVIDIA DOCA Telemetry Service (DTS), and Fluent Bit script are running on the target GPU servers. If a required service is not running, the script blocks the collection. For more information about GPU metrics, refer to [DCGM Exporter](#).
- On the machine where you are running the script, you will need:
 - Bash 4.0 or higher
 - Linux: Usually pre-installed
 - macOS: Install via Homebrew (see installation steps below)
 - SSH client
 - Standard Unix tools: tar, awk, sed, sha256sum/shasum
 - Internet connection for the initial download
- On the target GPU servers:
 - NVIDIA drivers installed (nvidia-smi available)
 - lspci (pciutils package)
 - ip command (iproute2 package)
 - rdma (iproute2 package)
 - ps (procps package)
 - SSH access (key-based or password authentication)

Download the GPU/NIC Topology Fleet Collector script

To download the GPU/NIC Topology Fleet Collector script:

1. Navigate to the **AI settings** page.

Edit <fabric> Settings > AI settings

2. Download the GPU/NIC Topology Fleet Collector script from the **AI settings** page.

Run the GPU/NIC Topology Fleet Collector script (single-server collection)

Test the script on a single-server collection before using on a fleet-wide deployment.

1. Verify that the DCGM Exporter is running on the GPU servers before running the GPU/NIC Topology Fleet Collector script.

See [DCGM Exporter](#) for more information.

2. Locate the GPU/NIC Topology Fleet Collector script that Cisco provided.
3. Run the GPU/NIC Topology Fleet Collector script on a single-server collection.
 - o On a single-server collection, run this command:

```
./gpu_nic_fleet_collector.sh > output.json
```

- o For debugging, run this command:

```
DEBUG=1 ./gpu_nic_fleet_collector.sh > output.json
```

- o For remote single host, use fleet mode:

```
./gpu_nic_fleet_collector.sh -f <(echo "gpu-server-01") -u root
```

Run the GPU/NIC Topology Fleet Collector script (fleet collection)

Once you have successfully tested the script on a single-server collection, use it on a fleet-wide deployment. Ensure that the DCGM Exporter, DTS, and Fluent Bit services are running before you execute the script. If a required service is not running, the script blocks the collection. To resolve this, start the required service and run the collection again.

1. Create a text file and add the GPU servers to this text file.

These GPU servers must be reachable through SSH.

For example:

- a. On a Windows system, right-click on the screen and choose **New > Text Document**.
- b. Add the GPU servers to this text document, with each GPU server on its own line. For example:

```
gpu-compute1
gpu-compute2
gpu-compute3
gpu-compute4
```

- c. Save and close the text file.
2. Verify your configuration without running the script.
 - o Linux:

```
bash ./gpu_nic_fleet_collector.sh --dry-run -f hosts.txt -u root
```

- o macOS:

```
/opt/homebrew/bin/bash ./gpu_nic_fleet_collector.sh --dry-run -f hosts.txt -u root
```

3. Run the collection.

- o Basic collection (sequential):

- Linux:

```
bash ./gpu_nic_fleet_collector.sh -f hosts.txt -u root
```

- macOS:

```
/opt/homebrew/bin/bash ./gpu_nic_fleet_collector.sh -f hosts.txt -u root
```

- o Recommended: Parallel collection (much faster):

- Linux:

```
bash ./gpu_nic_fleet_collector.sh -f hosts.txt -u root -j 10
```

- macOS:

```
/opt/homebrew/bin/bash ./gpu_nic_fleet_collector.sh -f hosts.txt -u root -j 10
```

You should see output similar to this:

```
=== GPU/NIC Topology Fleet Collection ===
```

```
Timestamp: 20260128_143022
```

```
Output: /path/to/out_20260128_143022
```

```
Hosts: 5 (local: no)
```

```
Parallelism: 10
```

```
[1/5] gpu-server-01 ... ✓
```

```
[2/5] gpu-server-02 ... ✓
```

```
[3/5] gpu-server-03 ... ✓
```

```
[4/5] 192.168.1.100 ... ✓
```

```
[5/5] 192.168.1.101 ... ✓
```

```
Collection complete!
```

```
Bundle: out_20260128_143022/fleet_bundle_20260128_143022.tar.gz
```

4. Review the output.

After the collection completes, you'll have this structure:

```
out_20260128_143022/
```

```
└─ fleet_bundle_20260128_143022.tar.gz    # Compressed bundle (deliverable)
```

```

├─ fleet_bundle_20260128_143022.tar.gz.sha256 # Checksum for verification

├─ mappings/

|   ├─ gpu-server-01_server_mapping.json    # Per-host inventory
|   └─ gpu-server-02_server_mapping.json

|   └─ ...

├─ logs/

|   ├─ gpu-server-01.log                    # Execution logs (for troubleshooting)
|   └─ ...

└─ metadata/

├─ MANIFEST.txt                            # Bundle metadata

├─ SHA256SUMS.txt                          # File checksums

└─ gpu_nic_fleet_collector.sh              # Script copy

```

5. Upload the zip file to **Job monitoring telemetry**.

This is the zip file that you will use in the **Job monitoring telemetry** area in the **AI settings** when you edit fabric settings for the AI fabric. See [AI settings](#) for more information.



You might also want to keep the zip file locally for troubleshooting.

Collect data in secured mode

Data collection on Nexus Dashboard from AI servers is supported in **http** and **https** modes. When you choose the **https** mode, mutual TLS (mTLS) is used and both the server and Nexus Dashboard to verify each other. This section provides the necessary steps to set up Nexus Dashboard to collect data in the **https** mode.

In order to accomplish this, you will need these certificates:

- **Server side:** Use the server certificate that you used to start the DCGM exporter. See [DCGM Exporter](#) for more information.
- **Client side:** On the Nexus Dashboard side:
 - CA certificate: Use the server certificate that you used to start the DCGM exporter. See [DCGM Exporter](#) for more information.

- o Client certificate and key



When using these procedures, we assume that the DCGM Exporter is also started in secure mode with the server certificate and key. See [DCGM Exporter](#) for more information.

Configure Nexus Dashboard AI settings to upload mapping file and start collection in secure mode

1. Upload the certificates to Nexus Dashboard.

See [Managing Certificates](#) for more information.

- a. Navigate to **Admin > Authentication and more > Certificate Management**.
- b. Upload the CA certificate.
 - i. Click **CA certificates**.
 - ii. Click **Actions > Add CA certificate**.
- c. Upload the system certificate, if necessary.

This will be the system certificate that you will use to connect to the DCGM Exporter.

- i. Click **System certificates**.
 - ii. Click **Actions > Add system certificate**.
2. Attach the system certificate to the feature.
 - a. In the **System certificates** area, choose the system certificate that you want to use to connect to the DCGM Exporter, then click **Actions > Manage feature attachment**.
 - b. In the **Manage feature attachments** page, click on the drop-down list and choose **AIComputeVisibility** from the list.
 - c. Click **Save**.



You can attach only one certificate to the **AIComputeVisibility** feature, which can be used on all the servers.

3. Upload the mapping file.
 - a. Navigate to **Manage > Fabrics**.
 - b. Click on the AI fabric.
 - c. Click **Actions > Edit fabric settings**.
 - d. Click **AI settings**, then scroll down to the **Job monitoring telemetry** area.
 - e. Upload the mapping file collected from the servers using the fleet collector script.

See [AI settings](#) and [GPU/NIC Topology Fleet Collector](#) for more information.

4. Add the system certificate in the **AI settings** page.
 - a. In the **Telemetry data pull mode** field, click **Secured SSL**.
 - b. In the **Client certificate name** field, enter the name of the system certificate that you attached

to the feature in step 2.

c. Click **Save**.



After you click **Save**, this configuration will be accepted only if the certificate in the **Client certification name** field has been uploaded and attached to the **AIComputeVisibility** feature.

Regardless of whether you chose the **Secured SSL** or the **Unsecured** option in the **Telemetry data pull mode** field, after you click **Save**, Nexus Dashboard starts collecting GPU data from the DCGM Exporter, such as memory, power consumption, and temperature.

Associate a fabric with an AI cluster

An AI cluster is essentially a fabric group. You can associate a fabric with an AI cluster if you see the **Associate with AI cluster** option in that fabric, as described in the steps below.

You must associate at least one fabric with an AI cluster in order to view the information on AI resources, as described in [View information on AI resources](#). If you do not have an AI cluster already created, you will create one as part of the process of associating a fabric with the AI cluster, as described below.

1. Click **Manage > Fabrics**.

2. In the **Fabrics** page, click the fabric in the **Name** column that you want to associate with an AI cluster.

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

3. Click **Actions > Configuration > Associate with AI cluster**.

The **Associate fabric with AI cluster** page appears.

4. Click the dropdown list in the **Associate fabric with AI cluster** page.

- o If you have an AI cluster already created and you want to associate the fabric with that AI cluster, choose it from the dropdown list.

- o If you want to create a new AI cluster, either because you do not have an AI cluster already created or you want to create a new AI cluster for this fabric:

- a. Click + **Create new AI cluster**.

The **Create new AI cluster** page appears.

- b. Enter a name for the new AI cluster, then click **Save**.

The new AI cluster appears in the **Associate fabric with AI cluster** page.

5. Click **Save** to associate the fabric with this AI cluster.

If you want to disassociate a fabric from an AI cluster:

1. Click **Manage > Fabrics**.

2. In the **Fabrics** page, click the fabric in the **Name** column that you want to disassociate from an AI cluster.

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

3. Click **Actions > Configuration > Remove from AI cluster**.

The AI cluster is also removed when you remove the last remaining fabric that is associated with it.

Associate a switch with a scalable unit

A scalable unit refers to a foundational, repeatable building block of GPU servers typically connected to the same set of leaf switches. These scalable units are connected using spine switches to deliver specific cluster sizes of AI compute capacity.

Key aspects of a scalable unit include:

- It consists of a set of network leaf switches and the GPU compute nodes connected to them.
- Different scalable unit types exist, defined by the Nexus leaf switch model and the number of UCS nodes, which directly impact the GPU cluster size.
- Scalable units come in various sizes. For example:
 - 8x Cisco N9332D-GX2B switches with up to 16 Cisco UCS C885A nodes supporting up to 128 GPUs
 - 8x Cisco N9364D-GX2A switches with up to 32 Cisco UCS C885A nodes supporting up to 256 GPUs
 - 8x Cisco N9364E-SG2 switches with up to 32 Cisco UCS C880A nodes supporting up to 256 GPUs
- Larger AI clusters are built by scaling out these scalable units in a leaf-spine design, ensuring predictable and consistent performance.
- The networks of scalable units can be managed through Nexus Dashboard, providing operational consistency across fabrics.

To associate a switch with a scalable unit:

1. Click **Manage > Fabrics**.
2. In the **Fabrics** page, click the fabric in the **Name** column that you want to associate with an AI cluster.

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

3. Click **Inventory**.

The **Inventory** page appears, with the **Switches** tab selected and configured switches displayed by default.

4. Click the check boxes next to two leaf switches that you want to associate with the scalable unit.
5. Click **Actions > Configuration > Associate with scalable unit**.

The **Associate with scalable unit** page appears.

6. Click the dropdown list in the **Associate with scalable unit** page.

- o If you have scalable unit already created and you want to associate the leaf switches with that scalable unit, choose it from the dropdown list.
- o If you want to create a new scalable unit, either because you do not have a scalable unit already created or you want to create a new scalable unit:
 - a. Click + **Create new scalable unit**.

The **Create new scalable unit** page appears.

- b. Enter a name for the new scalable unit, then click **Create**.

The new scalable unit appears in the **Associate with scalable unit** page.

7. Click **Save** to associate the leaf switches with this scalable unit.

If you want to disassociate leaf switches from a scalable unit:

1. Click **Manage > Fabrics**.
2. In the **Fabrics** page, click the fabric in the **Name** column where you want to disassociate switches from a scalable unit.

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

3. Click **Inventory**.

The **Inventory** page appears, with the **Switches** tab selected and configured switches displayed by default.

4. Click the check boxes next to the leaf switches that you want to disassociate from the scalable unit.
5. Click **Actions > Configuration > Disassociate with scalable unit**.

The scalable unit is also removed when you remove the last remaining leaf switch that is associated with it.

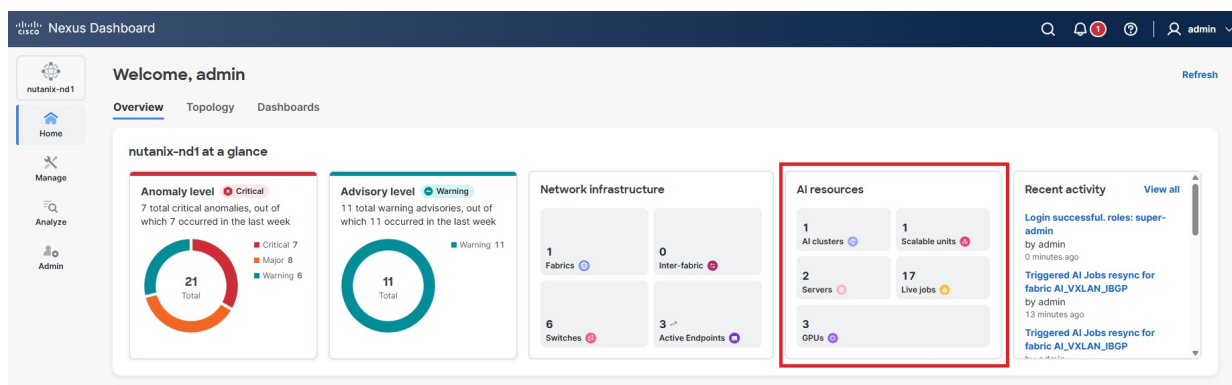
View information on AI resources



You must associate at least one fabric with an AI cluster in order to view the information on AI resources as described below. See [Associate a fabric with an AI cluster](#) for more information.

Information on AI resources is provided in two locations:

- A brief summary of the AI resources is provided in the **Overview** page (navigate to **Home > Overview**, then locate the **AI resources** tile), which displays this summary information:
 - o **AI clusters**: Displays the number of AI clusters managed by this system.
 - o **Scalable units**: Displays the number of scalable units.
 - o **Servers**: Displays the number of servers in the AI fabrics.
 - o **GPUs**: Displays the number of GPUs.
 - o **Jobs**: Displays the number of active AI jobs running in the clusters this system manages.



Click on any area in the **AI resources** tile to bring up a window with more detailed information on that area.

- More detailed information on AI resources is provided in the **Inventory** page (navigate to **Manage > Inventory > AI resources**).

These sections describe the more detailed information provided under **AI resources** in the **Inventory** page.

- [AI clusters](#)
- [Scalable units](#)
- [Servers](#)
- [GPUs](#)
- [Jobs](#)

AI clusters



You must associate at least one fabric with an AI cluster in order to view the information on AI clusters as described below. See [Associate a fabric with an AI cluster](#) for more information.

1. Navigate to **Manage > Inventory > AI resources**.
2. Click the **AI clusters** tab to view AI resource information at the cluster level.

Name	Member fabrics
AI_CLUSTER_SANJOSE_BUILDING20_3rd_FLOOR	AI_VXLAN_IBGP

The **AI clusters** tab displays this information:

Field	Description
Name	Displays the name of the AI cluster.
Member fabrics	Displays the fabrics that are members of each AI cluster. Click on a fabric name to navigate to the overview page for that fabric.

Scalable units



You must associate at least two leaf switches with a scalable unit in order to view the information on scalable units as described below. See [Associate a switch with a scalable unit](#) for information.

1. Navigate to **Manage > Inventory > AI resources**.
2. Click the **Scalable units** tab to view AI resource information at the scalable units level.

Inventory						Refresh
Switches	AI resources					
AI clusters	Scalable units	Servers	GPUs	Jobs		
Filter by attributes						
Scalable unit	AI cluster	Fabric	Switches	Servers	GPUs	
SCALABLE_UNIT_SANJOSE_BUILDING	AI_CLUSTER_SANJOSE_BUILDING20	AI_VXLAN_IBGP	2	2	3	

The **Scalable units** tab displays this information:

Field	Description
Scalable units	Displays the name of the scalable unit.
Cluster	Displays the cluster that contains the scalable unit.
Fabric	Displays the fabric that contains the scalable unit.
Switches	Displays the number of switches that contains the scalable unit.
Servers	Displays the number of servers that contains the scalable unit.
GPUs	Displays the number of GPUs within the scalable unit.

Servers

When you create an AI fabric, Nexus Dashboard automatically enables the LLDP feature on the network switches and discovers the GPU servers. However, you must also enable LLDP on the GPU servers:

```
root@gpu-compute1:~# apt install lldpd
Reading package lists... Done
```

```
Building dependency tree... Done
Reading state information... Done
lldpd is already the newest version (1.0.18-1build3).
0 upgraded, 0 newly installed, 0 to remove and 8 not upgraded.
root@gpu-compute1:~#
```

1. Navigate to **Manage > Inventory > AI resources**.
2. Click the **Servers** tab to view AI resource information at the server level.

Inventory Refresh

Switches **AI resources**

AI clusters Scalable units **Servers** GPUs Jobs

Anomaly Level

1 Critical 0 Major 0 Minor 0 Warning 1 Healthy

Filter by attributes

Anomaly level	Server Name	AI cluster	Fabric	Scalable unit	Management IP Address	Vendor	GPUs
Critical	gpu-compute1	AI_CLUSTER_SANJOSE_BI	AI_VXLAN_IBGP	SCALABLE_UNIT_SANJOS		NVIDIA	2
Healthy	gpu-compute3	AI_CLUSTER_SANJOSE_BI	AI_VXLAN_IBGP	SCALABLE_UNIT_SANJOS		NVIDIA	1

The **Servers** tab displays this information:

Field	Description
Anomaly Level	Displays the number of anomalies and their levels for the servers: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy

You can also view this information on each of the servers:

Field	Description
Anomaly level	Displays the anomaly level of a specific server with one of these values: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy
Server name	Displays the server name.
Cluster	Displays the cluster that contains the server.
Fabric	Displays the fabric that contains the server.
Scalable unit	Displays the scalable unit that is assigned to the server.
Management address	IP Displays the management IP address of the server.
Vendor	Displays the server's vendor.
GPUs	Displays the number of GPUs in this server.

GPUs

1. Navigate to **Manage > Inventory > AI resources**.
2. Click the **GPUs** tab to view AI resource information at the GPU level.

The screenshot shows the 'Inventory' page with the 'AI resources' tab selected. Under 'GPUs', there is a summary bar with the following counts: 2 Critical (red icon), 0 Major (orange icon), 0 Minor (yellow icon), 0 Warning (teal icon), and 1 Healthy (green icon). Below this is a filter bar with 'Filter by attributes' and buttons for 'List view' and 'Health view'. The main table displays GPU information for the path 'AI_CLUSTER_SANJOSE_BUILDING20_3rd_FLOOR > AI_VXLAN_IBGP'. The table has columns for Anomaly level, GPU ID, Servers, Vendor, Model, and Switch interface. The data rows are as follows:

Anomaly level	GPU ID	Servers	Vendor	Model	Switch interface
Healthy	0	gpu-compute3	NVIDIA	H100 NVL	AIML-Leaf1-Ethernet1/6
Critical	0	gpu-compute1	NVIDIA	H100 NVL	AIML-Leaf1-Ethernet1/1
Critical	1	gpu-compute1	NVIDIA	H100 NVL	AIML-Leaf2-Ethernet1/1

The **GPUs** tab displays this information:

Field	Description
GPU anomaly levels	Displays the number of anomalies and their levels for the GPUs: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy
GPU	Displays critical issues for GPUs in a cluster. Expand the scalable unit to view the connected GPUs.
List view/Health view	Determine how you want the GPU information to be presented. ▪ To view GPU information in a list format, click List view. ▪ To view GPU information in a grid format, click Health view. In Health view, each box represents a server, and each red square is a GPU within the server. Hover over a red square to view general information about a specific GPU. Click the box with the arrow to view more detailed information about a specific GPU. In Health view, Nexus Dashboard displays GPUs in good condition with reduced brightness and displays GPUs with health issues at full brightness. A status legend is available from the GPU health title and includes these states: Empty GPU slot, Healthy, Warning, Minor, Major, and Critical. If a server has empty GPU slots, the view shows a separate icon for the empty slots. When you filter by a GPU status, the view shows servers that have at least one GPU in that state and keeps the other GPUs in those servers visible for context.
The following fields become available if you click List View above.	

Field	Description
AI scalable unit	Displays if there is a critical issue with one or more AI scalable units. Click the down arrow to expand the view to see the AI scalable units in the GPU, where this information is displayed: - Anomaly level: Displays the anomaly level for each GPU: - Critical - Major - Minor - Warning - Healthy - GPU ID: The ID of the GPU. Click the number to bring up a window with additional information on a specific GPU. - Servers: Displays the server name associated with the GPU. - Vendor: Displays the vendor for the GPU. - Model: Displays the model number for the GPU. - Switch interface: Displays the switch interface associated with the GPU.

Jobs

1. Navigate to **Manage > Inventory > AI resources**.
2. Click the **Jobs** tab to view AI resource information at the jobs level.

Inventory

Switches **AI resources** Refresh Re-sync

AI clusters Scalable units Servers GPUs **Jobs**

Last day All jobs

State

Running Suspended Pending 16 Completed Failed 1 Others

Anomaly Level

16 Critical Major Minor Warning Healthy

Filter by attributes

ID	Name	Fabric	Status	Anomaly level	GPUs	Servers	Partition	Start Time	Run time	Time limit	Completion time	User
816	load_job_comp1	AI_VXLAN_IBGP	Cancelled	Critical	2	1	gpu-own	Dec 22 2025, 14:03:46	2m55s	20m0s	Dec 22 2025, 14:06:41	root
815	load_job_comp1	AI_VXLAN_IBGP	Completed	Critical	2	1	gpu-own	Dec 22 2025, 12:31:49	12m3s	20m0s	Dec 22 2025, 12:43:52	root
814	load_job_comp1	AI_VXLAN_IBGP	Completed	Critical	2	1	gpu-own	Dec 22 2025, 12:19:47	12m2s	20m0s	Dec 22 2025, 12:31:49	root
813	multi_node_job	AI_VXLAN_IBGP	Completed	Critical	2	2	gpu-own	Dec 22 2025, 06:22:05	12m3s	20m0s	Dec 22 2025, 06:34:08	root
812	multi_node_job	AI_VXLAN_IBGP	Completed	Critical	2	2	gpu-own	Dec 22 2025, 05:51:15	12m2s	20m0s	Dec 22 2025, 06:03:17	root
811	multi_node_job	AI_VXLAN_IBGP	Completed	Critical	2	2	gpu-own	Dec 22 2025, 05:33:32	12m2s	20m0s	Dec 22 2025, 05:45:34	root
810	multi_node_job	AI_VXLAN_IBGP	Completed	Critical	2	2	gpu-own	Dec 22 2025, 05:15:07	12m3s	20m0s	Dec 22 2025, 05:27:10	root
809	multi_node_job	AI_VXLAN_IBGP	Completed	Critical	2	2	gpu-own	Dec 22 2025, 00:55:57	12m3s	20m0s	Dec 22 2025, 01:08:00	root
808	multi_node_job	AI_VXLAN_IBGP	Completed	Critical	2	2	gpu-own	Dec 22 2025, 00:38:11	12m3s	20m0s	Dec 22 2025, 00:50:14	root
807	multi_node_job	AI_VXLAN_IBGP	Completed	Critical	2	2	gpu-own	Dec 22 2025, 00:05:56	12m3s	20m0s	Dec 22 2025, 00:17:59	root

17 items found Rows per page 10 < 1 2 >

The **Jobs** tab displays this information:

Field	Description
State	Displays the state of the AI jobs, with the number of AI jobs that are in these states: ▪ Running ▪ Suspended ▪ Pending ▪ Completed ▪ Failed ▪ Others: Displayed for all others states apart from the above.
Anomaly level	Displays the anomaly level of the AI jobs, with the number of AI jobs that are at these anomaly levels: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy

You can also view this information on each of the AI jobs:

Field	Description
ID	Displays the ID and name of an AI job. Click the ID or name to bring up a window with this additional information on a specific AI job: ▪ Overview: Job level ▪ Topology: Job level ▪ Optics: Job level ▪ Interface analytics: Job level ▪ Anomalies: Job level ▪ Activity logs: Job level
Name	
Fabric	Displays the fabric where the AI job is running.
Status	Displays the status of a specific AI job with one of these values: ▪ Running ▪ Suspended ▪ Pending ▪ Completed ▪ Failed: An AI job might fail if the GPU resources are not sufficient to run that AI job.

Field	Description
Anomaly level	Displays the anomaly level of a specific AI job with one of these values: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy
GPUs	Displays the number of GPUs that are being used by an AI job. Click the number to bring up a window with additional information on the GPUs that are being used by this AI job.
Servers	Displays the number of servers associated with an AI job.
Partition	Displays the partition associated with an AI job.
Start time	Displays the time that the AI job started.
Run time	Displays the amount of time that the AI job has been running.
Time limit	Displays the time limit for the AI job.
Completion time	Displays the time that the AI job was completed.  If an AI job has a status of Running, then this field shows the estimated completion time for this AI job based on the time limit that was provided for this AI job.
User	Displays the user who initiated the AI job.

View AI job information

You can use the AI job dashboard to view details of jobs, anomalies, optics, congestion, drops, CRC errors, GPU details, and other job details such as completion time, start and finish time, and so on.

In Nexus Dashboard 4.3.1, the AI job dashboard also shows server NIC statistics for AI jobs when server NIC telemetry is available. You can use the dashboard to correlate a job with GPU health, switch-interface health, and server-interface health.

1. Determine whether you want to see AI job information for all of the fabrics in the system or for a specific fabric.
 - To see AI job information for all of the fabrics in the system, navigate to **Manage > Inventory > AI resources**, then click the **Jobs** tab. See [Jobs](#) for more information.
 - To see AI job information for a specific fabric, navigate to **Manage > Fabrics**, then click on the AI fabric and click the **AI jobs** tab.
2. Review the information provided in the AI jobs page.

You can view AI job information at these levels:

- [View AI job information at the fabric level](#)

- [View information on a specific AI job](#)

View AI job information at the fabric level

- [Jobs: Fabric level](#)
- [Optics: Fabric level](#)
- [Analytics: Fabric level](#)
- [Anomalies: Fabric level](#)
- [Activity logs: Fabric level](#)

Jobs: Fabric level

The **Jobs** tab displays information on all of the AI jobs associated with this fabric.

AI_VXLAN_IBGP

Refresh

View in topology

Actions

Overview

Inventory

Connectivity

AI jobs

Segmentation and security

Configuration policies

Anomalies

Advisories

Integrations

History

Jobs

Optics

Analytics

Anomalies

Activity logs

Last day

All jobs

State

0 Running

0 Suspended

0 Pending

16 Completed

0 Failed

1 Others

Anomaly Level

16 Critical

0 Major

0 Minor

0 Warning

1 Healthy

Filter by attributes

Re-sync

ID	Name	Status	Anomaly level	GPUs	Servers	Start Time	Run time	Time limit	Completion time	User
816	load_job_comp1	Cancelled	Critical	2	1	Dec 22 2025, 14:03:46	2m55s	20m0s	Dec 22 2025, 14:06:41	root
815	load_job_comp1	Completed	Critical	2	1	Dec 22 2025, 12:31:49	12m3s	20m0s	Dec 22 2025, 12:43:52	root
814	load_job_comp1	Completed	Critical	2	1	Dec 22 2025, 12:19:47	12m2s	20m0s	Dec 22 2025, 12:31:49	root
813	multi_node_job	Completed	Critical	2	2	Dec 22 2025, 06:22:05	12m3s	20m0s	Dec 22 2025, 06:34:08	root
812	multi_node_job	Completed	Critical	2	2	Dec 22 2025, 05:51:15	12m2s	20m0s	Dec 22 2025, 06:03:17	root
811	multi_node_job	Completed	Critical	2	2	Dec 22 2025, 05:33:32	12m2s	20m0s	Dec 22 2025, 05:45:34	root
810	multi_node_job	Completed	Critical	2	2	Dec 22 2025, 05:15:07	12m3s	20m0s	Dec 22 2025, 05:27:10	root
809	multi_node_job	Completed	Critical	2	2	Dec 22 2025, 00:55:57	12m3s	20m0s	Dec 22 2025, 01:08:00	root
808	multi_node_job	Completed	Critical	2	2	Dec 22 2025, 00:38:11	12m3s	20m0s	Dec 22 2025, 00:50:14	root
807	multi_node_job	Completed	Critical	2	2	Dec 22 2025, 00:05:56	12m3s	20m0s	Dec 22 2025, 00:17:59	root

17 items found

Rows per page

10

<

1

2

>

To change the frequency and type of information displayed in the **Jobs** tab:

- Click the dropdown on **Last day** to change the frequency that the job information is updated.
- Click the dropdown on **All jobs** to change the type of job information displayed.

The **Jobs** tab displays this information:

Field	Description
State	Displays the state of the AI jobs, with the number of AI jobs that are in these states: ▪ Running ▪ Suspended ▪ Pending ▪ Completed ▪ Failed: An AI job might fail if the GPU resources are not sufficient to run that AI job.
Anomaly level	Displays the anomaly level of the AI jobs, with the number of AI jobs that are at these anomaly levels: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy

You can also view this information on each of the AI jobs:

Field	Description
ID	Displays the ID and name of an AI job.
Name	Click the ID or name to bring up a page with additional information on a specific AI job. See View information on a specific AI job for more information.
Status	Displays the status of a specific AI job with one of these values: ▪ Running ▪ Suspended ▪ Pending ▪ Completed ▪ Failed
Anomaly level	Displays the anomaly level of a specific AI job with one of these values: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy

Field	Description
GPUs	Displays the number of GPUs that are being used by an AI job. Click the number to bring up a page with additional information on the GPUs that are being used by this AI job.
Servers	Displays the number of servers associated with an AI job.
Start time	Displays the time that the AI job started.
Run time	Displays the amount of time that the AI job has been running.
Time limit	Displays the time limit for the AI job.
Completion time	Displays the time that the AI job was completed.  If an AI job has a status of Running, then this field shows the estimated completion time for this AI job based on the time limit that was provided for this AI job.
User	Displays the user who initiated the AI job.

Optics: Fabric level

The **Optics** tab displays optics (switch interface) information for AI jobs associated with this fabric.

AI_VXLAN_IBGP
Refresh
View in topology
Actions

Overview
Inventory
Connectivity
AI jobs
Segmentation and security
Configuration policies
Anomalies
Advisories
Integrations
History

Jobs
Optics
Analytics
Anomalies
Activity logs

3 Total server-facing transceivers
0 Critical
0 Major
0 Minor
0 Warning
3 Healthy

Filter by attributes

Switch interface	Metric	Anomaly	Neighbor-interface	Detection time
A1ML-Leaf1-eth1/1	CURRENT	Healthy	gpu-compute1-5000.e645.1a80	-
A1ML-Leaf1-eth1/1	RECEIVE_POWER	Healthy	gpu-compute1-5000.e645.1a80	-
A1ML-Leaf1-eth1/1	TRANSMIT_POWER	Healthy	gpu-compute1-5000.e645.1a80	-
A1ML-Leaf1-eth1/1	VOLTAGE	Healthy	gpu-compute1-5000.e645.1a80	-
A1ML-Leaf1-eth1/1	TEMPERATURE	Healthy	gpu-compute1-5000.e645.1a80	-
A1ML-Leaf1-eth1/6	CURRENT	Healthy	gpu-compute3-e89e.499d.408c	-
A1ML-Leaf1-eth1/6	RECEIVE_POWER	Healthy	gpu-compute3-e89e.499d.408c	-
A1ML-Leaf1-eth1/6	TRANSMIT_POWER	Healthy	gpu-compute3-e89e.499d.408c	-
A1ML-Leaf1-eth1/6	VOLTAGE	Healthy	gpu-compute3-e89e.499d.408c	-
A1ML-Leaf1-eth1/6	TEMPERATURE	Healthy	gpu-compute3-e89e.499d.408c	-

15 items found

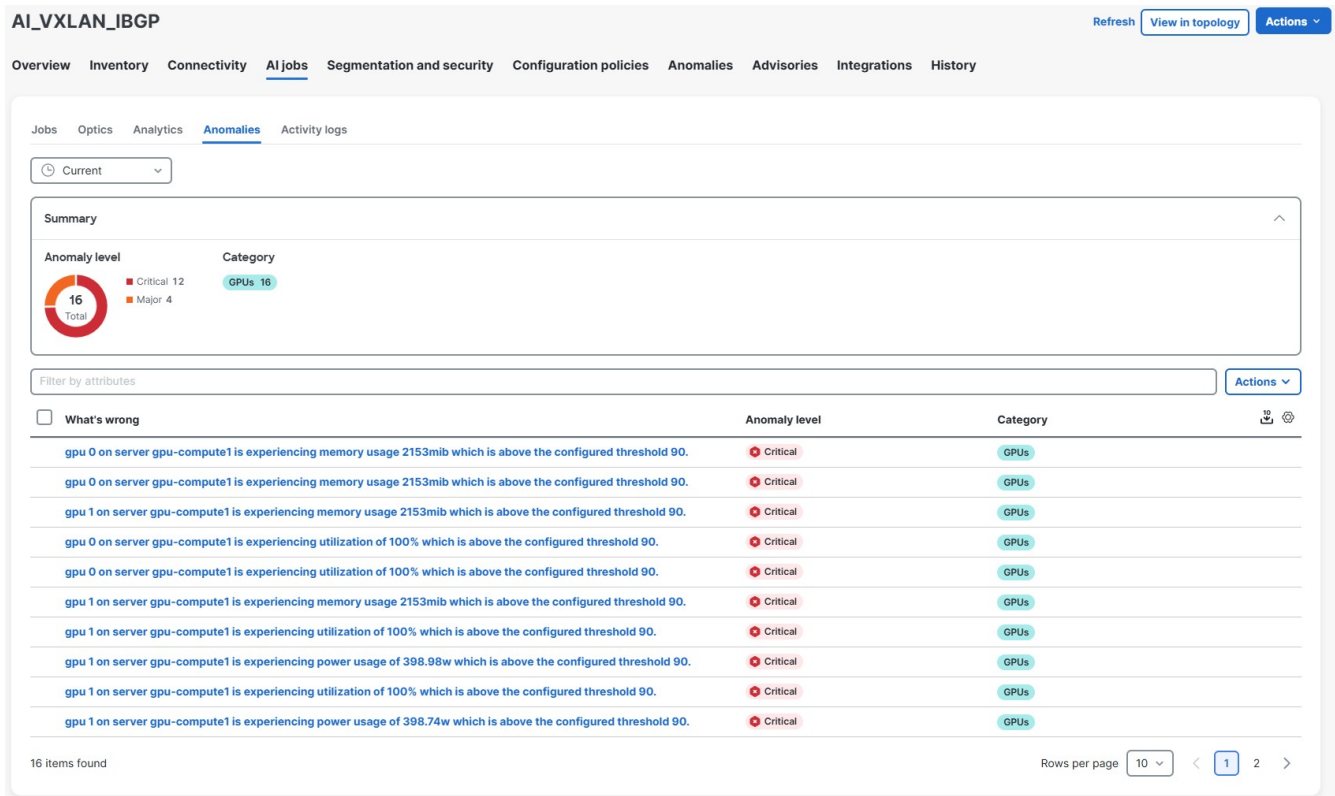
Rows per page
10
1
2

Field	Description
Switch interface	Displays the names of the switch interfaces. Click the name to bring up a page with this additional information on a specific switch interface.
Metric	Displays the metric for each switch interface.

Field	Description
Power Rx	Shows the received (Rx) or transmitted (Tx) power for AI jobs in the fabric. The Max value that exceeded threshold (dbm) field shows the power value. The Thresholds area shows two sets of low and high thresholds: * Alarm thresholds: Define the high and low values for a critical status. * Warning thresholds: Define the high and low values for a warning status.
Power Tx	
Voltage	Shows the voltage information for for AI jobs associated with this fabric. The Max value that exceeded threshold field shows the voltage value, and the Thresholds area shows the low and high thresholds for the voltage. A value above the minimum threshold is considered a warning, and a value above the high threshold is considered critical.
Current	Shows the current information for for AI jobs associated with this fabric. The Max value that exceeded threshold field shows the current value, and the Thresholds area shows the low and high thresholds for the current. A value above the minimum threshold is considered a warning, and a value above the high threshold is considered critical.
Temperature	Shows the temperature information for for AI jobs associated with this fabric. The Max value that exceeded threshold © field shows the temperature value, and the Thresholds area shows the low and high thresholds for the temperature. A value above the minimum threshold is considered a warning, and a value above the high threshold is considered critical.

Anomalies: Fabric level

The **Anomalies** tab displays anomaly information for AI jobs associated with this fabric.



To change the frequency for the information that is displayed in the **Anomalies** tab, click the dropdown on **Last day** to change the frequency that the anomaly information is updated:

- Current
- Last 15 minutes
- Last hour (default)
- Last 2 hours
- Last 24 hours
- Last week
- Last month
- Custom range

The **Anomalies** tab displays this information:

Field	Description
Summary	Provides a summary of the number of anomalies, and their levels and categories, found with the AI jobs in this fabric.
What's wrong	Provides a description of the anomaly that was found.

Field	Description
Anomaly level	Displays the anomaly level of the AI job at these anomaly levels: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy
Category	Displays the of anomalies grouped by various categories, such as Hardware, Capacity, Compliance, Connectivity, Configuration, Integrations, Active bugs, Telemetry configuration, GPUs, and System.

Activity logs: Fabric level

The **Activity logs** tab displays activity log information for AI jobs associated with this fabric.

AI_VXLAN_IBGP

Refresh

View in topology

Actions

OverviewInventoryConnectivityAI jobsSegmentation and securityConfiguration policiesAnomaliesAdvisoriesIntegrationsHistory

JobsOpticsAnalyticsAnomaliesActivity logs

Filter by attributes

ID	Status	Job name	Timestamp (PST)	Description
816	Cancelled	load_job_comp1	Dec 22 2025, 14:06:41	Job cancelled by user root
816	Running	load_job_comp1	Dec 22 2025, 14:03:46	Job started running
815	Completed	load_job_comp1	Dec 22 2025, 12:43:52	Job completed successfully
815	Running	load_job_comp1	Dec 22 2025, 12:31:49	Job started running
815	Pending	load_job_comp1	Dec 22 2025, 12:20:49	Job submitted by user root
814	Completed	load_job_comp1	Dec 22 2025, 12:31:49	Job completed successfully
814	Running	load_job_comp1	Dec 22 2025, 12:19:47	Job started running
813	Completed	multi_node_job	Dec 22 2025, 06:34:08	Job completed successfully
813	Running	multi_node_job	Dec 22 2025, 06:22:05	Job started running
812	Completed	multi_node_job	Dec 22 2025, 06:03:17	Job completed successfully

37 items found

Rows per page

10

<1234>

Use the time-range selector and the filter by attributes field to narrow the activity log entries that are displayed for the fabric.

The **Activity logs** tab displays this information:

Field	Description
ID	The ID of the activity log.
Status	The status of the activity log: ▪ Completed ▪ Running
Job name	The job that is associated with the activity log. Click the entry in the job name field to bring up information on that AI job. See View information on a specific AI job for more information.

Field	Description
Timestamp	The timestamp on the activity log.
Description	The description of the activity log, such as Job started running or Job completed successfully .
User	The user of the activity log.

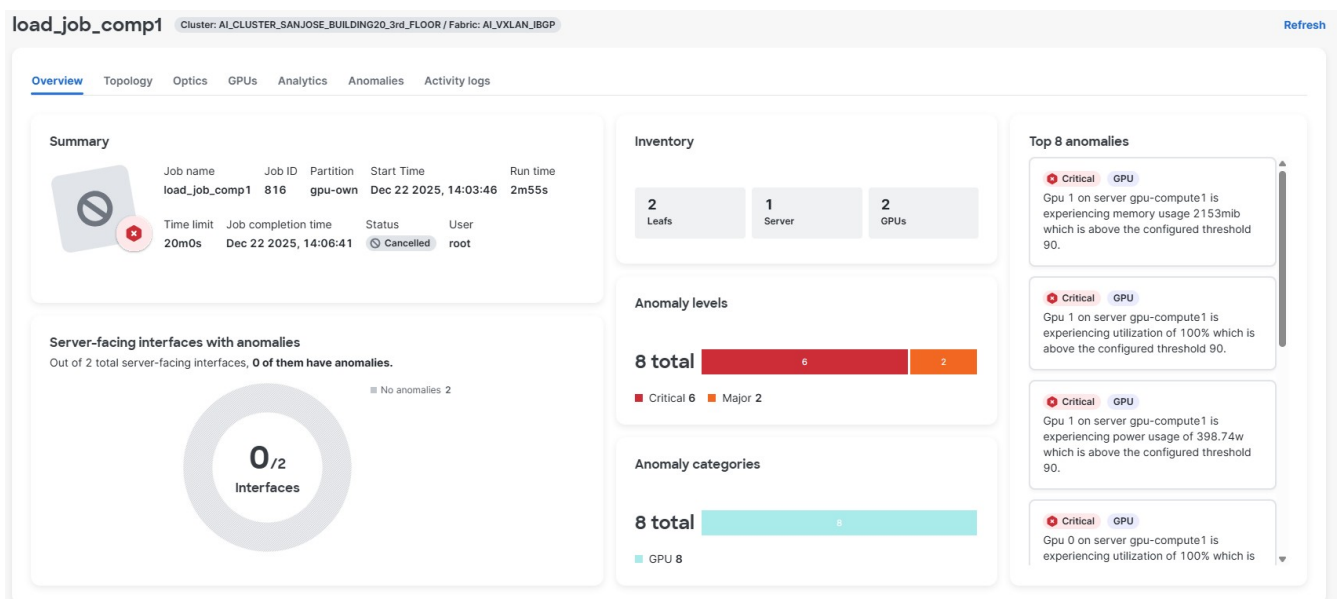
View information on a specific AI job

To view information on a specific AI job:

- Navigate to the page that displays all the AI jobs at either the system level or at a fabric level.
 - To see AI job information for all of the fabrics in the system, navigate to **Manage > Inventory > AI resources**, then click the **Jobs** tab.
 - To see AI job information for a specific fabric, navigate to **Manage > Fabrics**, then click on the AI fabric and click the **AI jobs** tab.
- Click on the **Jobs** tab, then click the ID or name of a specific job to bring up a page with this additional information on a specific AI job:
 - [Overview: Job level](#)
 - [Topology: Job level](#)
 - [Optics: Job level](#)
 - [GPUs: Job level](#)
 - [Interface analytics: Job level](#)
 - [Anomalies: Job level](#)
 - [Activity logs: Job level](#)

Overview: Job level

The **Overview** tab displays information as it relates to a specific job.



Field	Description
Summary	Provides a summary of the information available for this AI job, such as the job name and ID, the start and completion times for the job, the status of the job, and so on.
Inventory	Displays this information on the AI job: ▪ Leaf: The number of leaf switches associated with this AI job. ▪ Servers: The number of servers associated with this AI job. ▪ GPUs: The number of GPUs that are being used by this AI job.
Server-facing interfaces with anomalies	Provides information on the number of server-facing interfaces associated with this AI job and how many of the interfaces have anomalies.
Anomaly levels	Displays the number of anomalies, split by level, for the AI job.
Anomaly categories	Displays the number of anomalies by category for the AI job.
Anomalies	Displays the anomalies associated with the AI job.

Topology: Job level

The **Topology** tab displays topology information, similar to the standard **Topology** feature as described in [View LAN fabric components in Topology](#). However, this **Topology** tab provides centralized visualization and management of components as it relates to a specific AI job. For more information, see [AI endpoint and topology discovery](#).

The screenshot displays the 'load_job_comp1' interface for the cluster 'AL_CLUSTER_SANJOSE_BUILDING20_3rd_FLOOR / Fabric: ALVXLAN_IBGP'. The 'Topology' tab is active, showing a network diagram with components: 'Switches' (4), 'AIML-Leaf1', 'AIML-Leaf2', and 'gpu-compute1' (2). The overview pane on the right provides job details and anomalies.

Overview		
2 Leaves	1 Servers	2 GPUs
Job ID 816	Status Cancelled	Start Time Dec 22 2025, 14:03:46
Run time 2m55s	Time limit 20m0s	User root
Job completion time Dec 22 2025, 14:06:41		

Anomalies (8) [View anomalies](#)

- Critical GPU**
gpu 1 on server gpu-compute1 is experiencing memory usage 2153mib which is above the configured threshold 90.
- Critical GPU**
gpu 1 on server gpu-compute1 is experiencing utilization of 100% which is above the configured threshold 90.

The **Topology** tab includes an overview pane with job details, inventory counts, and a Critical anomalies summary for the selected job. Click **View anomalies** in the overview pane to view the anomalies that are associated with the job.

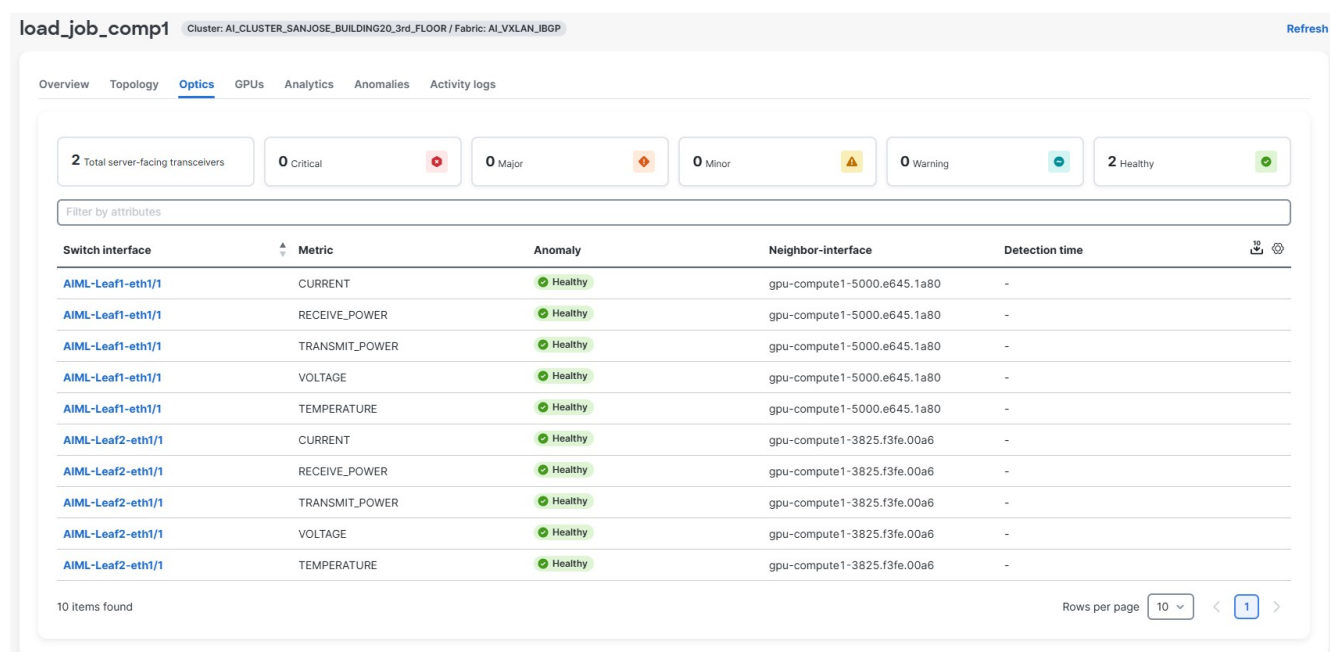


The AI job **Topology** view does not retain the topology from the time the job ran.

Instead, it displays the job topology using the currently available topology data and resource associations. If an AI cluster, fabric association, scalable unit, switch, server, NIC, or GPU changes after the job runs, the topology for that job might be incomplete.

Optics: Job level

The **Optics** tab displays information as it relates to a specific job.



Field	Description
Total server-facing interfaces with anomalies	Provides information on the number of server-facing interfaces associated with this AI job.
Interface anomaly levels	Displays the number of anomalies and their levels for the interfaces: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy

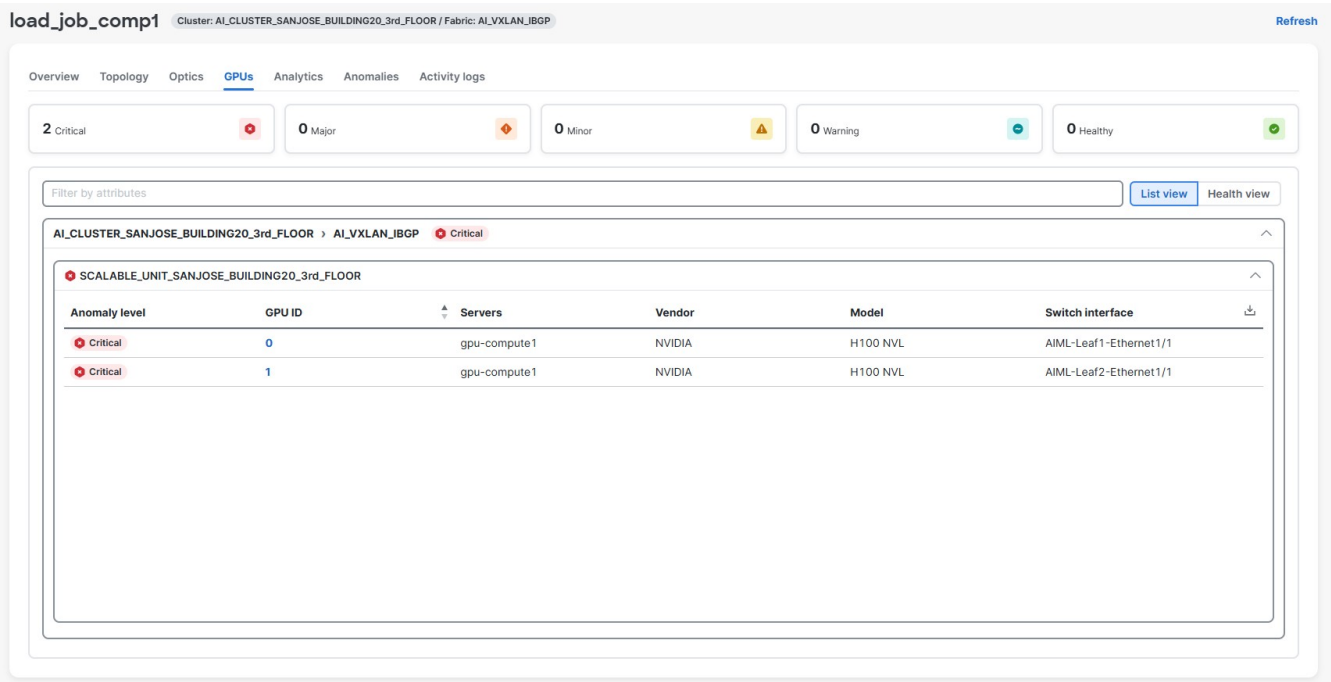
You can also view this information on each of the switch interfaces:

Field	Description
Switch interface	Displays the name of the switch interface. Click the name to bring up a page with this additional information on a specific switch interface.
Metric	Displays the metric for the switch interface.

Field	Description
Anomaly	Displays the anomaly level of a specific switch interface with one of these values: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy
Neighbor-interface	Displays the neighbor interface for the switch interface along with the neighbor device hostname and its MAC-based identifier.
Detection time	Displays the detection time.

GPUs: Job level

The **GPUs** tab displays information as it relates to a specific job.

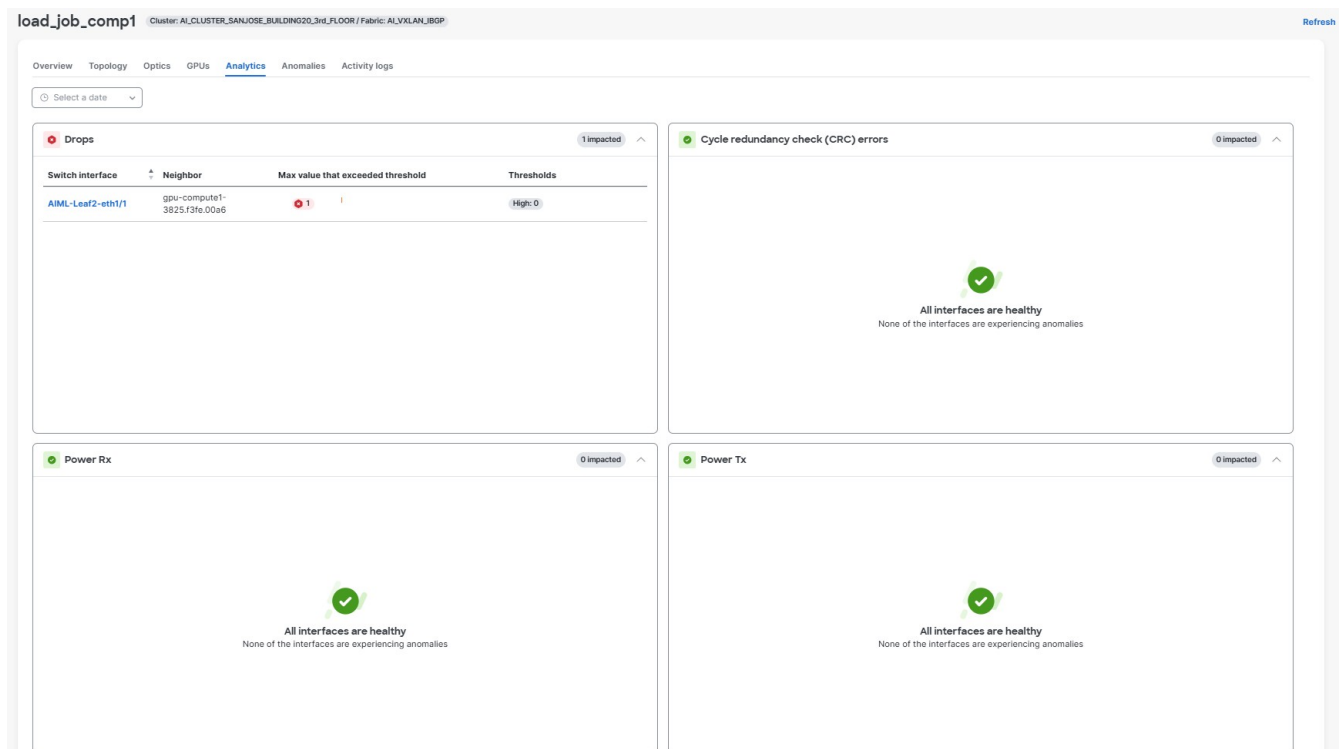


Field	Description
Anomaly level	Displays the number of anomalies and their levels for the GPUs: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy

Field	Description
GPU ID	Displays the ID of the GPU in the scalable unit. Click the link to open a new page, which provides details on this GPU, such as which job is running on it, time graph for memory usage, power usage, temperature and utilization.
Servers	Displays the server name associated with the GPU.
Vendor	Displays the vendor for the GPU.
Model	Displays the model number for the GPU.
Switch interface	Displays the switch interface associated with the GPU. The GPUs tab includes a GPU health summary, a filter, and List view and Health view options. In Health view, each box represents a server, and each square is a GPU within the server. Hover over a square to view general information about a specific GPU. Click the box with the arrow to view more detailed information about a specific GPU.

Interface analytics: Job level

The information that is provided in the **Analytics** tab is refreshed every 30 minutes or so. Click **Refresh** to refresh the information immediately.



The **Interface analytics** tab also displays server-interface metrics that are correlated with the selected AI job. Server-interface metrics appear in sections such as **Errors** and **Congestion**. Server-interface widgets can include **RDMA Nack - Rx**, **RDMA Nack - Tx**, **CNP Tx**, and **CNP Rx**. Switch-interface widgets continue to show switch-side interface, optics, utilization, and congestion information when switch telemetry is available.

By default, each analytics widget shows the top 10 interfaces for the selected time range. Click **Load all** interfaces to retrieve all available interfaces. Click a server-interface row to open the server

interface details page. Click a switch-interface row to open the switch interface details page.

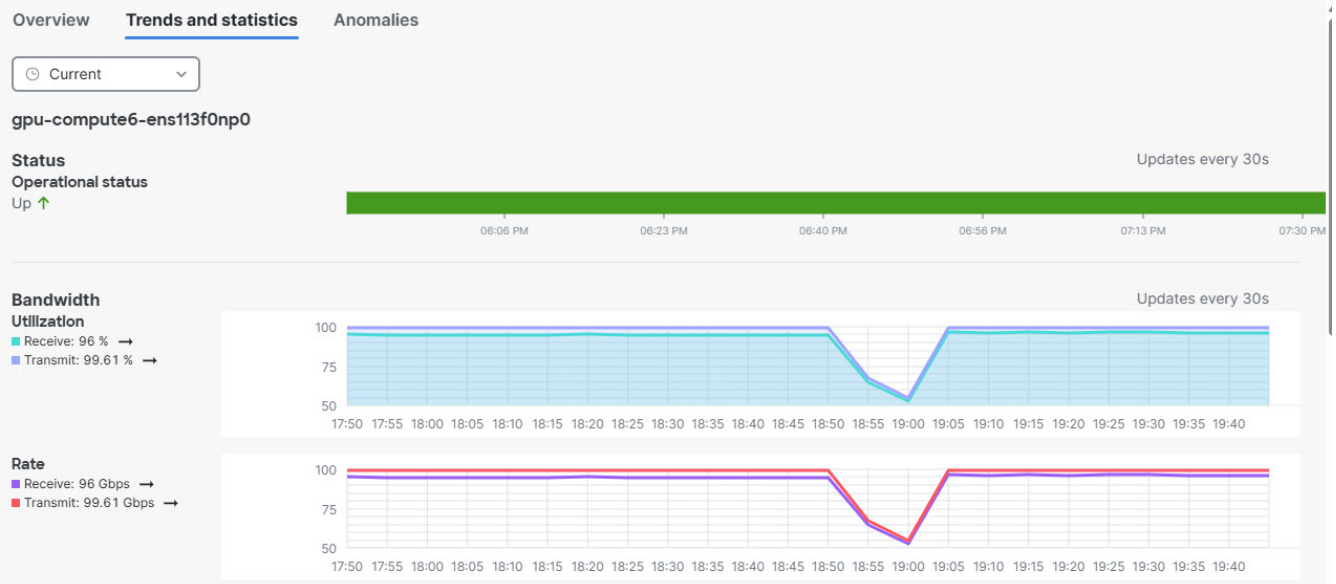
The **Analytics** tab displays this information.

Field	Description
Drops	Shows the number of drops that occurred for this job.
Cycle redundancy check (CRC) errors	Shows the number of cycle redundancy error checks that occurred for this job.
Power Rx	Shows the power Rx or TX information for this job.
Power Tx	The Max value that exceeded threshold (dbm) field shows the power Rx or TX value, and the Thresholds area shows the low and high thresholds for the power Rx or TX. A value above the minimum threshold is considered a warning, and a value above the high threshold is considered critical.
Voltage	Shows the voltage information for this job.
	The Max value that exceeded threshold field shows the voltage value, and the Thresholds area shows the low and high thresholds for the voltage. A value above the minimum threshold is considered a warning, and a value above the high threshold is considered critical.
Current	Shows the current information for this job.
	The Max value that exceeded threshold field shows the current value, and the Thresholds area shows the low and high thresholds for the current. A value above the minimum threshold is considered a warning, and a value above the high threshold is considered critical.
Module temperature	Shows the temperature information for this job.
	The Max value that exceeded threshold °C field shows the temperature value, and the Thresholds area shows the low and high thresholds for the temperature. A value above the minimum threshold is considered a warning, and a value above the high threshold is considered critical.

Server interface details page

Click a **Server interface** to open the *Server Interface Details* page. This page shows details for the selected server NIC interface.

Server Interface Detail for gpu-compute6-ens113f0np0



The server interface page includes these tabs:

- **Overview:** Displays the Anomaly level, General, and artificial intelligence (AI) jobs cards. The Anomaly level card shows the current anomaly status and count. The General card shows interface, NIC, operations speed, IP addresses, operational status, neighbor name, neighbor interface, and neighbor type. The AI jobs card shows jobs that use the server interface, including job name, start time, run time, time limit, job ID, job completion time, and user.
- **Trends and statistics:** Displays time-series data for server NIC status, bandwidth, errors, and congestion. Server NIC statistics in this tab include cyclic redundancy check (CRC) errors, transmit (Tx) discards, received (Rx) discards, bandwidth Rx, bandwidth Tx, negative acknowledgment (NACK) Rx, NACK Tx, congestion notification packet (CNP) Rx, CNP Tx, operational status, and link speed. When the Current toggle is on, the charts display the latest data for the selected interface.
- **Anomalies:** Displays anomalies that are associated with the selected server interface.

The **Overview** tab displays this information.

Field	Description
Anomaly level	Displays the highest anomaly severity and the total number of anomalies for the selected interface.
Interface	Displays the name of the server NIC interface.
NIC	Displays the type of NIC associated with the interface.
Operational speed	Displays the current speed of the interface.
IP addresses	Displays the IP addresses assigned to the interface.
Operational status	Displays whether the interface is operationally up or down.
Neighbor name	Displays the name of the connected neighbor device.
Neighbor interface	Displays the interface name of the connected neighbor device.
Neighbor type	Displays the type of connected neighbor device (for example, switch).

Field	Description
AI jobs	Displays details for AI jobs using this interface, including job name, start time, run time, time limit, job ID, job completion time, and user.

The **Trends and statistics** tab displays this information.

Field	Description
Status	Displays the current operational status of the interface.
Bandwidth Utilization	Displays the receive (Rx) and transmit (Tx) bandwidth utilization as a percentage.
Rate	Displays the receive (Rx) and transmit (Tx) data rates in gigabits per second (Gbps).
Cycle redundancy check (CRC) errors	Displays the number of CRC errors on the interface.
Transmit discard	Displays the number of packets discarded during transmission.
Receive discard	Displays the number of packets discarded during reception.
RDMA NACK - Rx/Tx	Displays the number of RDMA Negative Acknowledgment (NACK) events received (Rx) or transmitted (Tx).
Congestion notification packets (Tx/Rx)	Displays the number of congestion notification packets transmitted (Tx) or received (Rx).
Pause frames (Rx/Tx)	Displays the number of pause frames received (Rx) or transmitted (Tx).
Congestion indicator	Displays the congestion score for the interface.

The **Anomalies** tab displays this information.

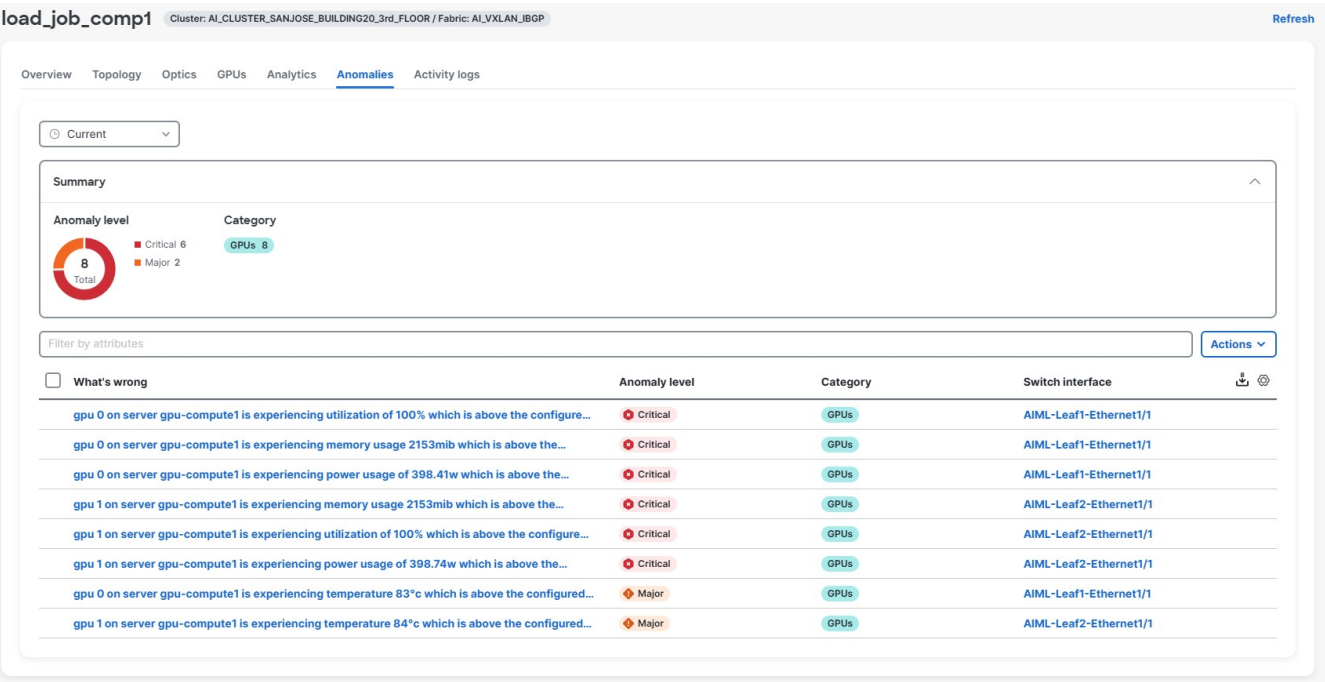
Field	Description
Anomaly type	Displays the anomaly name, such as Interface Ingress Error or Interface Utilization High Threshold.
Anomaly level	Displays the severity of the anomaly.
Category	Displays the anomaly category, such as Connectivity.
Total count	Displays the number of anomalies in the anomaly group.

Anomalies: Job level

To select the time window for which Nexus Dashboard displays **Anomalies**, click the dropdown on **Last day**.

- Current
- Last 15 minutes
- Last hour (default)
- Last 2 hours
- Last 24 hours

- Last week
- Last month
- Custom range



The **Anomalies** tab displays this information.

Field	Description
What's wrong	Provides a description of the anomaly that was found.
Anomaly level	Displays the anomaly level of the AI job at these anomaly levels: ▪ Critical ▪ Major ▪ Minor ▪ Warning ▪ Healthy

Field	Description
Category	Displays a subset of anomalies that are relevant in the context of AI jobs: ▪ Interface Dom Transmit Power ▪ Interface Dom Receive Power ▪ Interface Dom Current ▪ Interface Dom Temperature ▪ Interface Dom Voltage ▪ Interface FCS Error ▪ Interface Stomped CRC Error ▪ Interface Down ▪ Interface Utilization High Threshold ▪ Network Congestion Indication ▪ Environmental Outbound Discard High
Switch interface	Displays the name of the switch interface. Click the name to bring up a page with this additional information on a specific switch interface.

Server NIC anomalies can also appear in the **Anomalies** tab for metrics such as CRC errors, transmit discards, received discards, NACK Rx, and NACK Tx.

For server NIC anomalies, the **What's wrong** column identifies the affected server interface and describes the issue. Click the description to open the anomaly details page. When available, the **Switch interface** column shows the neighboring switch interface that is associated with the affected server interface.

Activity logs: Job level

The **Activity logs** page provides information on the timeline of the different states a jobs goes through, such as **Pending**, **Running**, **Completed**, and so on.

load_job_comp1
Cluster: AI_CLUSTER_SANJOSE_BUILDING20_3rd_FLOOR / Fabric: AI_VXLAN_IBGP
Refresh

Overview
Topology
Optics
GPUs
Analytics
Anomalies
Activity logs

Filter by attributes

ID	Status	Timestamp (PST)	Description	User
816	Cancelled	Dec 22 2025, 14:06:41	Job cancelled by user root	root
816	Running	Dec 22 2025, 14:03:46	Job started running	root

Use the **Filter by attributes** field to narrow the activity log entries that are displayed for the selected job.

The **Activity logs** tab displays this information:

Field	Description
ID	Displays the ID of AI job.
Status	Displays the status of the activity log: ▪ Completed ▪ Running
Timestamp (PST)	Displays the timestamp on the activity log.
Description	Displays the description of the activity log, such as Job started running or Job completed successfully .
User	Displays the user name of the activity log.

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