



BGP Fabric, Release 12.2.2/12.2.3

Table of Contents

New and Changed Information	1
Creating VXLAN EVPN Fabric with eBGP-based Underlay	2
Guidelines for VXLAN Fabric With eBGP Underlay	2
General Parameters	3
EVPN	4
vPC	9
Protocols	10
Security	12
Advanced	13
Freeform	16
Manageability	17
Bootstrap	17
Configuration Backup	19
Flow Monitor	20
Adding Switches	22
Assigning Switch Roles	23
AI/ML QoS Classification and Queuing Policies	24
About AI/ML QoS Classification and Queuing Policies	24
Guidelines and Limitations: AI/ML QoS Classification and Queuing Policies	25
Configuring AI/ML QoS Classification and Queuing Policies	25
Using the Custom QoS Templates to Create a Policy	26
Border Gateway Support	28
Border Gateway Route Filtering	28
Guidelines and Limitations	30
Layer 3 VNI Without VLAN	31
Guidelines and Limitations: Layer 3 VNI Without VLAN	31
MACsec Support in Data Center VXLAN EVPN and BGP Fabrics	33
Guidelines	33
Enabling MACsec	33
Disabling MACsec	38
vPC Fabric Peering	40
Guidelines and Limitations	40
QoS for Fabric vPC-Peering	41
Creating a Virtual Peer Link	42
Converting a Physical Peer Link to a Virtual Peer Link	43
Before you begin	43
Converting a Virtual Peer Link to a Physical Peer Link	45
Before you begin	45
Deploying Fabric Underlay eBGP Policies	46
Deploying Fabric Overlay eBGP Policies	47
Deploying Spine Switch Overlay Policies	47

Deploying Leaf Switch Overlay Policies	47
Adding a Super Spine Switch to an Existing VXLAN BGP EVPN Fabric	48
Copyright	50

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 nor of the new features up to this release.

Release Version	Feature	Description
NDFC 12.2.3 release	Support for Cisco encryption type 6 in BGP Authentication Key Encryption Type field	When creating a BGP fabric type, support is now available for Cisco encryption type 6 in the BGP Authentication Key Encryption Type field. See [Creating VXLAN EVPN Fabric with eBGPbased Underlay] for more information.
NDFC 12.2.3 release	Introduction of Freeform tab	When creating a BGP fabric type, the freeform configuration fields are now consolidated under a single Freeform tab. See [Creating VXLAN EVPN Fabric with eBGPbased Underlay] for more information.
NDFC 12.2.3 release	Support for QoS queuing policies for an interface speed of 800 Gb	When enabling the AI / ML QOS & Queuing Policy option, a new AI_Fabric_QOS_800G field is added to enable QoS queuing policies for an interface speed of 800 Gb. See Advanced for more information.
NDFC 12.2.2 release	Support added for creating VXLAN EVPN fabrics with a PIMv6 underlay and TRMv6	In previous releases of NDFC, NDFC supported an IPv6 underlay with ingress replication (IR). Beginning with the NDFC 12.2.2 release, NDFC added support for multicast replication. Previously NDFC supported a standalone VXLAN IPv4 fabric. Beginning with NDFC 12.2.2, NDFC supports creating a Multi-Site Domain (MSD) fabric with VXLAN IPv6. Prior to NDFC 12.2.2, NDFC supported Tenant Routed Multicast (TRM) IPv4. With NDFC 12.2.2, NDFC added support for TRMv6. With this release, new fields are added to the EVPN tab for configuring TRMv6. This feature is available for the following fabric types: • Data Center VXLAN EVPN fabric • BGP (eBGP EVPN) fabric • VXLAN EVPN Multi-Site fabric For more information, see EVPN and the section "Configuring VXLAN EVPN Fabrics with a PIMv6 Underlay and TRMv6" in Data Center VXLAN EVPN.

Creating VXLAN EVPN Fabric with eBGP-based Underlay

Guidelines for VXLAN Fabric With eBGP Underlay

- Brownfield migration is not supported for eBGP fabrics.
- You cannot change the leaf switch Autonomous System (AS) number after it is created and the configuration is deployed. You must delete the **leaf_bgp_asn** policy and perform **Recalculate & Deploy** to remove BGP configuration related to this AS. Then, add the **leaf_bgp_asn** policy with the new AS number.
- To switch between Multi-AS and Same-Tier-AS modes, remove all manually added BGP policies (including **leaf_bgp_asn** on the leaf switch and the EBGp overlay policies), and perform the **Recalculate & Deploy** operation before the mode change.
- You cannot change or delete the leaf switch **leaf_bgp_asn** policy if there are ebgp overlay policies present on the device. You need to delete the eBGP overlay policy first, and then delete the **leaf_bgp_asn** policy.
- The roles supported with VXLAN BGP fabric vary, depending on the release:
 - For releases prior to NDFC release 12.2.1, the roles supported with VXLAN BGP fabric are leaf, spine, border, super spine and border super spine.
 - For NDFC release 12.2.1 and later, the border gateway, border gateway spine, and border gateway super spine roles are also supported with VXLAN BGP fabric.
- Intra-fabric links only support IPv6 link local addresses when the underlay is IPv6.
- On a border device, VRF-Lite is supported with manual mode.
- Support for VXLAN eBGP fabrics in VXLAN multi-site varies, depending on the release:
 - For releases prior to NDFC release 12.2.1, VXLAN multi-site is not supported for VXLAN eBGP fabrics.
 - For release NDFC 12.2.1 and later, VXLAN eBGP fabrics are supported as child fabrics in VXLAN Multi-Site fabrics.
- TRM (Tenant Routed Multicast) is supported with an eBGP fabric with an IPv4 or an IPv6 underlay.
- VXLAN with an IPv6 underlay does not support the following features:
 - Bidirectional Forwarding Detection (BFD)
 - MACSec
 - Flexible Netflow
 - BGP authentication

To create a VXLAN EVPN fabric with an IPv4 or an IPv6 eBGP underlay:

1. Navigate to **Manage > Fabrics**.
2. From the **Actions** drop-down list, click **Create Fabric**.
3. Enter a unique name for the fabric in the **Fabric Name** field, then click **Choose Fabric**.

A list of all available fabric templates are listed.

4. From the available list of fabric templates, choose the **BGP Fabric** template, then click **Select**.

The fabric settings for creating a standalone fabric appear. Most of the fields are pre-filled for the fabric.

5. Enter the necessary field values or edit pre-filled fields, as required.

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

- [General Parameters](#)
- [EVPN](#)
- [vPC](#)
- [Protocols](#)
- [Security](#)
- [Advanced](#)
- [Freeform](#) (NDFC 12.2.3 release only)
- [Manageability](#)
- [Bootstrap](#)
- [Configuration Backup](#)
- [Flow Monitor](#)

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

- Click on the fabric to display a summary in the slide-in pane.
- Click on the **Launch** icon to display the **Fabric Overview** page.

General Parameters

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

General Parameters for VXLAN EVPN Fabric with eBGP

Field	Description
BGP ASN for Spines	Enter the autonomous system number (ASN) for the fabric's spine switches.
BGP ASN for Super Spines	Enter the ASN used for super spine and border super spines, if the fabric contains any super spine or border super spine switches.
BGP AS Mode	Choose Multi-AS or Same-Tier-AS. ▪ In a Multi-AS fabric, a unique AS number per leaf/border is used. ▪ In a Same-Tier-AS fabric, all leaf nodes share one unique AS and all border nodes share another unique AS. In both Multi-AS and Same-Tier-AS, all spine switches in a fabric share one unique AS number. The fabric is identified by the spine switch ASN.

Field	Description
Allow Same ASN On Leafs	Uses the same ASN on all the leaf nodes even when you have configured Multi-AS mode.
Enable IPv6 routed fabric or VXLAN with IPv6 underlay	Enables IPv6 routed fabric or IPv6 underlay. With the checkbox cleared, the system configures IPv4 routed fabric or IPv4 underlay. To configure IPv6 underlay, you must also configure the VXLAN overlay parameters in the EVPN tab.
Underlay Subnet IP Mask	Specifies the subnet mask for the fabric interface IP addresses.
Manual Underlay IP Address Allocation	Check the check box to disable dynamic underlay IP address allocations.
Underlay Routing Loopback IP Range	Specifies the loopback IPv4 addresses for protocol peering.
Underlay Subnet IP Range	Specifies the IPv4 addresses for underlay P2P routing traffic between interfaces.
Underlay Routing Loopback IPv6 Range	Specifies the loopback IPv6 addresses for protocol peering.
Disable Route-Map Tag	Disables subnet redistribution.
Route-Map Tag	Configures a route tag for redistributing subnets. By default, the tag value of 12345 is configured, when enabled.
Subinterface Dot1q Range	Specifies the subinterface range when Layer 3 sub interfaces are used.
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.  Performance monitoring is supported on switches with NX-OS Release 9.3.6 and later.

EVPN

EVPN configuration for VXLAN EVPN Fabric with eBGP

Field	Description
Enable EVPN VXLAN Overlay	Enables the VXLAN overlay provisioning for the fabric. You can convert a routed fabric to a VXLAN enabled fabric by selecting this option. When the fabric is VXLAN enabled, you can create and deploy overlay networks or VRF instances. The procedure for creating and deploying networks or VRFs is the same as in Data Center VXLAN EVPN. For more information, see the "Creating Network for Standalone Fabrics" and "Creating VRF" sections in Data Center VXLAN EVPN. You must uncheck the Enable EVPN VXLAN Overlay check box to create a routed fabric (an IP fabric with no VXLAN encapsulation). In a routed fabric, you can create and deploy networks. For more information, see the section "Overview of Networks in a Routed Fabric" in Managing BGP-Based Routed Fabrics. Whether you create an eBGP routed or eBGP VXLAN fabric, the fabric uses eBGP as the control plane to build intra-fabric connectivity. Links between spine and leaf switches are auto-configured with point-to-point (P2P) numbered IP addresses with eBGP peering built on top. If a network or a VRF is created in a fabric, you cannot switch between VXLAN EVPN mode and routed fabric mode by selecting the Enable EVPN VXLAN Overlay check box. You need to delete these networks or VRFs to change the fabric setting. Routed_Network_Universal Template is applicable to a routed fabric only. When you convert the routed fabric to a VXLAN EVPN fabric, set the network template and network extension template to the ones defined for VXLAN EVPN: Default_Network_Universal and Default_Network_Universal. If you have a customized template for a VXLAN EVPN fabric, you can also choose to use it.
First Hop Redundancy Protocol	This field is available if you did not select the Enable EVPN VXLAN Overlay option above. This field is only applicable to a routed fabric. Specify the First Hop Redundancy Protocol that you want to use: • hsrp: Hot Standby Router Protocol • vrrp: Virtual Router Redundancy Protocol  After a network has been created, you cannot change this fabric setting. You should delete all networks, and then change the First Hop Redundancy Protocol setting.
The following fields in the EVPN tab are only applicable if you selected the Enable EVPN VXLAN Overlay option above.	
Anycast Gateway MAC	Specifies the anycast gateway MAC address for the leaf switches.

Field	Description
Enable VXLAN OAM	Enables the VXLAN operations, administration, and maintenance (OAM) function for existing switches. This is enabled by default. Uncheck the check box to disable VXLAN OAM feature. If you want to enable the VXLAN OAM on specific switches and disable on other switches in the fabric, use freeform configurations to enable OAM and disable OAM in the fabric settings.  VXLAN OAM feature in Cisco NDFC is supported on a single fabric or site only. VXLAN OAM is not supported with multi-site fabrics.
Enable Tenant DHCP	Enables tenant DHCP support.
vPC advertise-pip	Check the check box to enable the advertise PIP (primary IP address) feature on vPC enabled leaf or border leaf switches.
vPC advertise-pip on Border only	Check the check box to enable advertise-pip on vPC border switches and border gateways only. Applicable only when the vPC advertise-pip option is not enabled.
Replication Mode	Specifies the mode of replication that is used in the fabric - ingress replication or multicast.
Multicast Group Subnet	Specifies the IP address prefix used for multicast communication. A unique IP address is allocated from this group for each overlay network. This field is used when the underlay is IPv4 and replication mode is multicast.
IPv6 Multicast Group Subnet	Specifies the IPv6 address prefix used for multicast communication. A unique IP address is allocated from this group for each overlay network. This field is used when the underlay is IPv6 and replication mode is multicast.
Enable IPv4 Tenant Routed Multicast	Check the check box to enable Tenant Routed Multicast (TRM) for IPv4 support for VXLAN EVPN fabrics. For more information, see the section "Creating a VXLAN EVPN Fabric Using the Data Center VXLAN EVPN template" in Data Center VXLAN EVPN .
Enable IPv6 Tenant Routed Multicast (TRMv6)	Check the check box to enable Tenant Routed Multicast (TRM) for IPv6 support for VXLAN EVPN fabrics. For more information, see the section "Configuring VXLAN EVPN Fabrics with a PIMv6 Underlay and TRMv6" in Data Center VXLAN EVPN .
Default MDT Address for TRM VRFs	Indicates the multicast address for TRMv4 traffic. 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	Indicates the multicast address for TRMv6 traffic. 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 the IPv6 Multicast Group Subnet field. For more information, see the section "Overview of Tenant Routed Multicast" in Configuring Tenant Routed Multicast.

Field	Description
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 enable multicast mode, only the fields pertaining to that multicast mode is enabled and the fields related to other the multicast mode are disabled.  BIDIR-PIM is supported on Cisco's Cloud Scale Family platforms 9300-EX and 9300-FX/FX2, and with NX-OS Release 9.2(1) and later.
Underlay RP Loopback ID	Specifies the loopback ID used for the Rendezvous Point (RP). The default is 254.
The following fields are enabled if you choose bidir as the RP mode. Depending on the RP count, either 2 or 4 phantom RP loopback ID fields are enabled.	
Underlay Primary RP Loopback ID	Specifies the primary loopback ID used for phantom RP.
Underlay Backup RP Loopback ID	Specifies the secondary (or backup) loopback ID used for Fallback Bidir-PIM phantom RP.
The following Loopback ID options are applicable only when the RP count is 4 and if bidir is chosen.	
Underlay Second Backup RP Loopback ID	Specifies the second backup loopback ID used for phantom RP.
Underlay Third Backup RP Loopback ID	Specifies the third backup loopback ID used for phantom RP.
MVPN VRI ID Range	Beginning with NDFC 12.2.2, use this field for allocating a unique MVPN VRI ID per vPC. This field is needed for the following conditions: ▪ 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. ▪ If you are using a switch running in Layer 3 VNI mode.  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.

Field	Description
Enable MVPN VRI ID Re-allocation	Beginning with NDFC release 12.2.2, enable this check box to generate a one-time VRI ID reallocation. NDFC 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.
VRF Template	Specifies the VRF template for creating VRFs, and the VRF extension template for enabling VRF extension to other fabrics.
VRF Extension Template	
Network Template	Specifies the network template for creating networks, and the network extension template for extending networks to other fabrics.
Network Extension Template	
Overlay Mode	Specify the VRF/network configuration using one of the following options: • config-profile • cli The default option is config-profile. 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.
Allow L3VNI w/o VLAN	Beginning with NDFC release 12.2.2, check the box to allow Layer 3 VNI configuration without having to configure a VLAN.
Enable L3VNI w/o VLAN	Beginning with NDFC release 12.2.1, 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 "Creating a VRF" section in About Fabric Overview for LAN Operational Mode Setups.

Field	Description
Site Id	Beginning with NDFC release 12.2.1, eBGP fabrics are now supported as child fabrics in a VXLAN EVPN Multi-Site fabric. For more information, see Border Gateway Support . Enter a value in this field for the ASN for the BGP fabric in the VXLAN EVPN Multi-Site fabric. Valid entries for this field range from 1-281474976710655. If left empty, the value for this field defaults to the fabric ASN.
Anycast Border Gateway advertise-pip	Beginning with NDFC release 12.2.1, check the check box to advertise the Anycast Border Gateway PIP as a Virtual Tunnel End Point (VTEP) when adding a BGP fabric as a child fabric in a VXLAN EVPN Multi-Site fabric. This setting takes affect when you click Recalculate Config as part of the VXLAN EVPN Multi-Site fabric configuration process. For more information, see Border Gateway Support .
Underlay VTEP Loopback IP Range	Specifies the loopback IP address range for VTEPs. This is typically the Loopback1 IP address range.
Underlay RP Loopback IP Range	Specifies anycast or phantom RP IP address range.
Underlay VTEP Loopback IPv6 Range	Specifies the loopback IPv6 address range for VTEPs. This is typically the Loopback1 and Anycast Loopback IPv6 address range.
Layer 2 VXLAN VNI Range	Specify the VXLAN VNI IDs for the fabric.
Layer 3 VXLAN VNI Range	
Network VLAN Range	VLAN ranges for the Layer 3 VRF and overlay network.
VRF VLAN Range	
VRF Lite Deployment	Specifies the VRF Lite method for extending inter fabric connections. Only Manual is supported.

vPC

vPC Configuration for VXLAN EVPN Fabric with eBGP

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.
vPC Peer Keep Alive option	From the drop-down list, select management or loopback . To use IP addresses assigned to the management port and the management VRF, select management . To use IP addresses assigned to loopback interfaces (in non-management VRF), select 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.

Field	Description
vPC Delay Restore Time	Specifies the vPC delay restore period in seconds.
vPC Peer Link Port Channel Number	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.
Fabric wide vPC Domain Id	Enables the usage of same vPC Domain Id on all vPC pairs in the fabric. 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. Otherwise unique vPC domain IDs are used (in increment of 1) for each vPC pair.
Enable Qos for Fabric vPC-Peering	Enables 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 spines.

Protocols

Protocol Configuration for VXLAN EVPN Fabric with eBGP

Field	Description
Routing Loopback Id	The loopback interface ID is populated as 0 by default. It is used as the BGP router ID.
VTEP Loopback Id	The loopback interface ID is populated as 1 and it is used for VTEP peering purposes.
BGP Maximum Paths	Specifies maximum number for BGP routes to be installed for same prefix on the switches for ECMP.
Enable BGP Authentication	Check the check box to enable BGP authentication. If you enable this field, the BGP Authentication Key Encryption Type and BGP Authentication Key fields are enabled.
BGP Authentication Key Encryption Type	Choose the encryption type: • 3 for the 3DES encryption type • 6 for the Cisco encryption type 6 (available only in NDFC 12.2.3 and later) • 7 for the Cisco encryption type 7

Field	Description
BGP Authentication Key	Enter the encrypted key based on the encryption type.  Plain-text passwords are not supported. Log on to the switch, retrieve the encrypted key. Enter the key in the BGP Authentication Key field. For more information, refer to the section "Retrieving the Encrypted BFD Authentication Key" in IPFM and Classic IPFM.
Enable PIM Hello Authentication	Enables the PIM hello authentication.
PIM Hello Authentication Key	Specifies the PIM hello authentication key.
Enable BFD	Check the Enable BFD 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. NDFC supports BFD within a fabric. 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 BFD configurations requires configurations to be deployed via the per switch freeform or per interface freeform policies. The following configuration is pushed after you enable BFD. <pre>'feature bfd'</pre> For NDFC with BFD-enabled, the following configurations are pushed on all the P2P fabric interfaces: <pre>no ip redirects no ipv6 redirects</pre> For information about BFD feature compatibility, refer your respective platform documentation and for information about the supported software versions, see <i>Cisco Nexus Dashboard Fabric Controller Compatibility Matrix</i>.
Enable BFD for BGP	Check the check box to enable BFD for the BGP neighbor. This option is disabled by default.
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 enabled.
BFD Authentication Key ID	Specifies the BFD authentication key ID for the interface authentication.

Field	Description
BFD Authentication Key	Specifies the BFD authentication key. For information about how to retrieve the BFD authentication parameters, see refer to the section "Retrieving the Encrypted BFD Authentication Key" in IPFM and Classic IPFM.

Security

Security Configuration for VXLAN EVPN Fabric with eBGP

Field	Description
Enable MACsec	Check the box to enable MACsec in the BGP fabric. For more information, see Enabling MACsec .
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 .

Field	Description
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.  The MACsec configuration is not deployed on the switches after the fabric deployment is complete. You need to enable MACsec on intra-fabric links to deploy the MACsec configuration on the switch.
MACsec Status Report Timer	Specifies the MACsec operational status periodic report timer in minutes.

Advanced

Advanced Configuration for VXLAN EVPN Fabric with eBGP

Field	Description
Intra Fabric Interface MTU	Specifies the MTU for the intra fabric interface. This value must be an even number.
Layer 2 Host Interface MTU	Specifies the MTU for the Layer 2 host interface. This value must be an even number.
Power Supply Mode	Choose the appropriate power supply mode.
CoPP Profile	From the drop-down list, select the appropriate Control Plane Policing (CoPP) profile policy for the fabric. By default, the strict is selected.
VTEP HoldDown Time	Specifies the NVE source interface hold down time.
VRF Lite Subnet IP Range	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.
VRF Lite Subnet Mask	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.

Field	Description
Enable CDP for Bootstrapped Switch	Check the check box to enable CDP for switches discovered using Bootstrap.
Enable NX-API	Check the check box to enable NX-API on HTTPS. This check box is checked by default.  - There are subtle behavior changes to the EPL feature when the NX-API option is enabled or disabled. See the section EPL Behavior with NX-API Configuration for more information. - Disabling the NX-API option will change the behavior for the Layer 4 to Layer 7 service configuration feature. See the Guidelines and Limitations for Layer 4 to Layer 7 Services for more information.
Enable NX-API on HTTP	Specifies enabling of NX-API on HTTP. Check Enable NX-API on HTTP and Enable NX-API check boxes 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 NDFC, 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 both Enable NX-API and Enable NX-API on HTTP check boxes, applications use HTTP.
Enable Strict Config Compliance	Enable the Strict Configuration Compliance feature by selecting this check box. For more information, see the section "Strict Configuration Compliance" in Configuration Compliance.
Enable AAA IP Authorization	Enables AAA IP authorization (make sure IP Authorization is enabled in the AAA Server).
Enable NDFC as Trap Host	Check the check box to enable NDFC as a trap host.
Enable TCAM Allocation	TCAM commands are automatically generated for VXLAN and vPC Fabric Peering when enabled.
Greenfield Cleanup Option	Enable the switch cleanup option for greenfield switches without performing a switch reload. This option is typically recommended only for the Data centers with the Cisco Nexus 9000v Switches.

Field	Description
Enable Default Queuing Policies	Check the 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 deploy the configuration. Pre-defined QoS configurations are included that can be used for various Cisco Nexus 9000 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 necessary configuration to the peer interface freeform block. Review the actual queuing policies by opening the policy file in the template editor. From Cisco NDFC Web UI, choose Operations > Template. Search for the queuing policies by the policy file name, for example, queuing_policy_default_8q_cloudscale. Choose the file and click the Modify/View template icon to edit the policy. See the <i>Cisco Nexus 9000 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 Nexus 9200 Series Switches and the Cisco Nexus 9000 Series Switches that ends 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 FEXs. You can change from the queuing_policy_default_4q_cloudscale policy to the queuing_policy_default_8q_cloudscale policy only when FEXs are offline.
N9K R-Series Platform Queuing Policy	Select 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.
Enable AI/ML QoS and Queuing Policies	Beginning with NDFC release 12.2.1, check the box to enable AI/ML QoS and queuing policies in the BGP fabric. For more information, see AI/ML QoS Classification and Queuing Policies.  This option is not available if you also enabled either of the following options: • Enable Qos for Fabric vPC-Peering option in the vPC tab • Enable Default Queuing Policies option in the Advanced tab

Field	Description
AI / ML QOS & Queuing Policy	This field is available if you checked the Enable AI/ML QOS and Queuing Policies option above. Beginning with NDFC release 12.2.1, choose the queuing policy from the drop-down list based on the predominant fabric link speed for certain switches in the fabric. For more information, see AI/ML QoS Classification and Queuing Policies. Options are: • AI_Fabric_QOS_800G: New field added in NDFC release 12.2.3. Enable QoS queuing policies for an interface speed of 800 Gb. • AI_Fabric_QOS_400G: Enable QoS queuing policies for an interface speed of 400 Gb. • AI_Fabric_QOS_100G: Enable QoS queuing policies for an interface speed of 100 Gb. • AI_Fabric_QOS_25G: Enable QoS queuing policies for an interface speed of 25 Gb.
Priority flow control watch-dog interval	This field is available if you checked the Enable AI/ML QOS and Queuing Policies option above. Beginning with NDFC release 12.2.2, when enabling the AI/ML feature, priority-flow-control watchdog-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>.
Enable Real Time Interface Statistics Collection	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).
	The following fields are available in the Advanced tab in NDFC release 12.2.2, but are moved to the Freeform tab in NDFC release 12.2.3.
Leaf Freeform Config	Adds CLIs that should be added to switches that have the Leaf , Border , and Border Gateway roles.
Spine Freeform Config	Adds CLIs that should be added to switches with a Spine , Border Spine , Border Gateway Spine , and Super Spine roles.
Intra-fabric Links Additional Config	Adds CLIs that should be added to the intra fabric links.

Freeform

This tab is introduced in NDFC release 12.2.3.

The fields in this tab are shown below. For more information, see [Enabling Freeform Configurations on](#)

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

Manageability

Manageability Parameters for VXLAN EVPN Fabric with eBGP

Field	Description
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	Specifies the comma-separated list of IP addresses (v4/v6) of the NTP server.
NTP Server VRFs	Specifies one VRF for all NTP servers or a comma-separated list of VRFs, one per NTP server.
Syslog Server IPs	Specifies the comma-separated list of IP addresses (v4/v6) IP address of the syslog servers, if used.
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 configs. If AAA configs are specified in the fabric settings, switch_freeform PTI with source as UNDERLAY_AAA and description as AAA Configurations will be created.

Bootstrap

Bootstrap Parameters for VXLAN EVPN Fabric with eBGP

Field	Description
Enable Bootstrap	Check the Enable Bootstrap check box to enable the bootstrap feature. 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 - Check the Local DHCP Server check box and enter details for the remaining mandatory fields.
Enable Local DHCP Server	Check the Enable Local DHCP Server check box to enable DHCP service on NDFC and initiate automatic IP address assignment. When you check this check box, the DHCP Scope Start Address and DHCP Scope End Address fields become editable. If you do not check this check box, NDFC 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 NDFC IPv6 POAP is not supported with Cisco Nexus 7000 Series Switches. Cisco Nexus 9000 and 3000 Series Switches support IPv6 POAP only when switches are either Layer2 adjacent (eth1 or out-of-band subnet must be a /64) or Layer3 adjacent residing in some IPv6 /64 subnet. Subnet prefixes other than /64 are not supported.
DHCP Scope Start Address	Specifies the first and last IP addresses of the IP address range. IPs from this scope are allocated to the switches during the POAP bootstrap process.
DHCP Scope End Address	
Switch Mgmt Default Gateway	Specifies the default gateway for the DHCP scope.
Switch Mgmt IP Subnet Prefix	Specifies the prefix length for DHCP scope.
DHCP scope and management default gateway IP address specification	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.
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	Check the check box to include AAA configs from the Manageability tab during device bootup.

Field	Description
Bootstrap Freeform Config	Enter additional commands as needed. For example, if you are using AAA or remote authentication related configurations, you need to add these configurations in this field to save the intent. After the devices boot up, they contain the intent defined in the Bootstrap Freeform Config field. Copy-paste the running-configuration 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 "Resolving Freeform Config Errors in Switches" in Enabling Freeform Configurations on Fabric Switches.
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

Configuration Backup

Configuration Backup Parameters for VXLAN EVPN Fabric with eBGP

Field	Description
Hourly Fabric Backup	Check the Hourly Fabric Backup check box to enable an hourly backup of fabric configurations and the intent. You can enable an hourly backup for fresh fabric configurations and the intent. If there is a configuration push in the previous hour, NDFC takes a backup. Intent refers to configurations that are saved in NDFC but yet to be provisioned on the switches.
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	Specifies the scheduled backup time in a 24-hour format. This field is enabled if you check the Scheduled Fabric Backup check box. 1. Select both the check boxes to enable both back up processes. The backup process is initiated after you click Save.  Hourly and scheduled backup processes happen only during the next periodic configuration compliance activity, and there can be a delay of up to an hour. 2. To trigger an immediate backup, do the following: a. Choose Overview > Topology. b. Click within the specific fabric box. The fabric topology screen comes up. c. From the Actions pane at the left part of the screen, click Re-Sync Fabric. You can also initiate the fabric backup in the fabric topology window. Click Backup Now in the Actions pane. 3. Click Save after filling and updating relevant information.

Flow Monitor

Configuration Parameters for VXLAN EVPN Fabric with eBGP

Field	Description
Enable Netflow	Check the Enable Netflow check box to enable Netflow on VTEPs for this Fabric. By default, Netflow is disabled.  When Netflow is enabled on the fabric, you can choose not to have netflow on a particular switch by having a fake 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 NDFC, see section "Netflow Support" in Understanding LAN Fabrics.

Field	Description
Netflow Exporter	To add Netflow exporters for receiving netflow data: 1. In the Netflow Exporter area, choose Actions > Add. The Add Item page appears. 2. In the Exporter Name field, enter a name of the exporter. 3. In the IP field, enter the IP address of the exporter. 4. In the VRF field, specify the VRF over which the exporter is routed. 5. In the Source Interface field, enter the source interface name. 6. In the UDP Port field, enter the UDP port number over which the netflow data is exported. 7. Click Save to configure the exporter.
Netflow Record	To add Netflow records: 1. In the Netflow Record area, choose Actions > Add to add one or more Netflow records. 2. In the Record Name field, enter a name for the record. 3. In the Record Template field, select the required templates. From 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. - o netflow_ipv4_record - to use the IPv4 record template. - o netflow_l2_record - to use the Layer 2 record template. 4. Check the Is Layer2 Record check box if the record is for Layer2 netflow. 5. Click Save to configure the report.
Netflow Monitor	To add Netflow monitors: 1. In the Netflow Monitor area, choose Actions > Add. 2. In the Monitor Name field, enter a name for the monitor. 3. In the Record Name field, enter a name for the record. 4. In the Exporter1 Name field, enter the name of the exporter for the netflow monitor. 5. (Optional) In the Exporter2 Name field, enter 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 the Netflow Record and Netflow Exporter configuration sections in the Flow Monitor tab. . Click Save to configure the flow monitor.

Adding Switches

Switch 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 "Adding Switches to a Fabric" in [Add Switches for LAN Operational Mode](#).

Assigning Switch Roles

To assign roles to switches on Nexus Dashboard Fabric Controller refer to the "Assigning Switch Roles" section in [Add Switches for LAN Operational Mode](#).

AI/ML QoS Classification and Queuing Policies

The following sections provide information about the AI/ML QoS classification and queuing policies feature, introduced in NDFC release 12.2.1:

- [About AI/ML QoS Classification and Queuing Policies](#)
- [Guidelines and Limitations: AI/ML QoS Classification and Queuing Policies](#)
- [Configuring AI/ML QoS Classification and Queuing Policies](#)
- [Using the Custom QoS Templates to Create a Policy](#)

About AI/ML QoS Classification and Queuing Policies

Beginning with NDFC release 12.2.1, support is available for configuring a low latency, high throughput and lossless fabric configuration that can be used for artificial intelligence (AI) and machine learning (ML) traffic.

This 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/ML QoS policy, NDFC 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/ML QoS feature on a VXLAN EVPN fabric, then the Network Virtual (NVE) interface will have the attached AI/ML QoS policies.

Use the following areas to enable this feature:

- When configuring a BGP fabric, new fields are available to enable the feature and to set the queuing policy parameters based on the interface speed. For more information, see [Advanced](#).
- You can also use the following AI/ML-specific switch templates to create custom device policies, which can be used on host interfaces:
 - **AI_Fabric_QoS_Classification_Custom**: An interface template that is available for applying a custom queuing policy to an interface.
 - **AI_Fabric_QoS_Queueing_Custom**: A switch template that is available for user-defined queuing policy configurations.

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

Beginning with NDFC release 12.2.2, when enabling the AI/ML 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/ML 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 Nexus 9000 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: AI/ML QoS Classification and Queuing Policies

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

- 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 Nexus 9300-FX2, 9300-FX3, 9300-GX, and 9300-GX2 series switches.
- This feature is not supported in fabrics with devices that are assigned with a ToR role.

Configuring AI/ML QoS Classification and Queuing Policies

Follow these steps to configure AI/ML QoS and queuing policies:

1. Enable AI/ML 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/ML 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 [Add Interfaces for LAN Operational Mode](#) for more information.

- a. Within a fabric where you enabled AI/ML 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/ML 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/ML QoS and queuing policies on this host interface.

For example, these policy templates contain the necessary AI/ML 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/ML QoS for this host interface.

This enables the QoS classification on this interface if AI/ML 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 [Using the Custom QoS Templates to Create a Policy](#), 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 NDFC will use the default QOS_CLASSIFICATION policy, if available.

- g. If you created a custom queuing policy using the procedures provided in [Using the Custom QoS Templates to Create a Policy](#), 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/ML QoS and queuing policy configurations for this host interface.

Using the Custom QoS Templates to Create a Policy

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

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

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

2. Click the **Policies** tab.
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.

6. Select the appropriate custom Classification or Queuing template from the list, then click **Select**.

The following templates are specific to the AI/ML QoS and queuing policies feature that is introduced in NDFC release 12.2.1. Use these templates to create policies that can be used on one or more host interfaces:

- **AI_Fabric_QOS_Classification_Custom**: An interface template that is available for applying a custom queuing policy to an interface.
- **AI_Fabric_QOS_Queueing_Custom**: A switch template that is available for user-defined queuing policy configurations.

7. 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.

Border Gateway Support

Beginning with NDFC release 12.2.1, the following border gateway roles are now supported in an eBGP fabric:

- Border gateway
- Border gateway spine
- Border gateway super spine

Support is also available for eBGP fabrics as child fabrics within a VXLAN EVPN Multi-Site fabric. Several new fields are available when creating a BGP fabric that allow you to enter an ASN and to advertise the Anycast Border Gateway PIP as a Virtual Tunnel End Point (VTEP) for the BGP fabric in the VXLAN EVPN Multi-Site fabric. For more information, see [EVPN](#).

The way that the BGP ASN is allocated for the border gateway will match the settings for the leaf switches and border switches. You also might have to make additional configurations, based on the device role that you assign to the switch:

- If you set the device role as **border gateway**, then you must manually add the **leaf_bgp_asn** policy to the device to specify the BGP ASN that will be used on the switch, following the same Multi-AS mode or Same-Tier-AS mode rule that you have set at the fabric level.
- If you set the device role as **border gateway spine**, then the device ASN is the same value as the value provided in the **BGP ASN for Spines** option at the fabric level.
- If you set the device role as **border gateway super spine**, then the device ASN is the same value as the value provided in the **BGP ASN for Super Spines** option at the fabric level.

Nexus Dashboard Fabric Controller will automatically generate the BGP underlay configuration after you click **Recalculate & Deploy** at the end of the configuration process. You will then need to manually add the eBGP overlay policies for the border gateways using the procedures provided in [Deploying Fabric Overlay eBGP Policies](#).

Border Gateway Route Filtering

Two prefix-list policies are created for each border gateway when you perform a **Recalculate & Deploy** at the VXLAN EVPN Multi-Site fabric level:

- **fabric-internal:** This policy is applied to an eBGP session toward fabric internal, and is used to avoid routes that are learned from DCI from being propagated to other nodes within the fabric.

For this policy, permit only the following:

- Routing loopback (typically loopback0) subnet
- VTEP loopback (typically loopback1) subnet
- Anycast RP loopback (typically loopback254) subnet
- Multisite loopback IP address

Following is a sample configuration for this policy:

```
ip prefix-list ebgp-fabric-to-internal seq 10 permit 8.2.0.0/22 eq 32
ip prefix-list ebgp-fabric-to-internal seq 20 permit 8.3.0.0/22 eq 32
ip prefix-list ebgp-fabric-to-internal seq 30 permit 8.254.254.0/24 eq 32
ip prefix-list ebgp-fabric-to-internal seq 40 permit 10.10.0.2/32
```

```
route-map ebgp-rmap-filter-to-internal permit 10
  match ip address prefix-list ebgp-fabric-to-internal
route-map ebgp-rmap-filter-to-internal deny 20
```

```
router bgp 65028
  neighbor 8.4.0.1
    address-family ipv4 unicast
      route-map ebgp-rmap-filter-to-internal out
```

- **fabric-to-dci:** This policy is applied to the VXLAN EVPN Multi-Site underlay eBGP session, and is used to avoid VXLAN fabric underlay routes from being propagated to external fabrics.

For this policy, permit only the BGW and the local loopback interface IP address of its neighbor BGW (loopback0/1/VPC VIP/multisite loopback).

Following is a sample configuration for this policy:

```
ip prefix-list ebgp-fabric-to-dci seq 10 permit 8.2.0.2/32
ip prefix-list ebgp-fabric-to-dci seq 20 permit 8.3.0.1/32
ip prefix-list ebgp-fabric-to-dci seq 30 permit 10.10.0.2/32
```

```
route-map ebgp-rmap-filter-to-dci permit 10
  match ip address prefix-list ebgp-fabric-to-dci
route-map ebgp-rmap-filter-to-dci deny 20
```

```
router bgp 65028
  neighbor 10.10.1.2
    address-family ipv4 unicast
      route-map ebgp-rmap-filter-to-dci out
      next-hop-self
```

You can edit either of these two prefix-list policies using the template [ipv4_prefix_list_internal](#) to add the custom prefix-list entries. For more information on adding or editing a policy, see the section "Adding a Policy" in [About Fabric Overview for LAN Operational Mode Setups](#).

Guidelines and Limitations

Following are the guidelines and limitations with the border gateway support in BGP fabrics:

- Border gateway support for an eBGP fabric is only available when EVPN is enabled, and when the underlay is IPv4.
- Tenant Routed Multicast (TRM) is also supported with border gateway roles in a BGP fabric.

Layer 3 VNI Without VLAN

Beginning with NDFC release 12.2.1, the Layer 3 VNI without VLAN feature is now supported with Nexus Dashboard Fabric Controller. With this feature, Layer 3 VNI configurations no longer require a VLAN per VRF.

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 will be 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 "Creating a VRF" section in [About Fabric Overview for LAN Operational Mode Setups](#) 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: Layer 3 VNI Without VLAN](#))

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

Guidelines and Limitations: 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 Nexus 9000 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 Nexus 9000 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 will be 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.
- If you upgrade from an earlier release to NDFC 12.2.1, already-configured VRFs and fabrics will retain their existing pre-12.2.1 settings where the Layer 3 VNI without VLAN feature is disabled (the **Enable L3VNI w/o VLAN** box is unchecked). Once you are fully upgraded to NDFC release 12.2.1, you can manually change these settings to enable the Layer 3 VNI without VLAN feature, if desired.

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.

Beginning with NDFC 12.2.2, NDFC added support for MACsec for inter-fabric links, DCI MACsec, with a QKD server to generate keys or using preshared keys. For more information, see [Connecting Two Fabrics with MACsec Using QKD](#).

MACsec is supported on switches with minimum Cisco NX-OS releases 7.0(3), 17(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 Nexus 9000 Series NX-OS Security Configuration Guide*.

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

Enabling MACsec

The process described in this section is for configuring MACsec for intra-fabric links for Data Center VXLAN EVPN and BGP fabrics.

For information on configuring data center interconnect (DCI) MACsec for inter-fabric links and using a quantum key distribution (QKD) server, see [Connecting Two Fabrics with MACsec Using QKD](#).

1. Navigate to **Manage > Fabrics**.
2. Click **Actions > Create** to create a new fabric or click **Actions > Edit Fabric** on an existing Data Center VXLAN EVPN fabric.
3. Click the **Security** tab 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, NDFC 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, NDFC 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-XPN-128 • GCM-AES-XPN-256 The default value is GCM-AES-XPN-256.
DCI MACsec Primary Key String	Specify a Cisco Type 7 encrypted octet string that is used for establishing the primary DCI MACsec session. For AES_256_CMAC, the key string length must be 130 and for AES_128_CMAC, the key string length must be 66. If these values are not specified correctly, an error is displayed when you save the fabric.  The default key lifetime is infinite.
DCI MACsec Primary Cryptographic Algorithm	Choose the cryptographic algorithm used for the primary key string. It can be AES_128_CMAC or AES_256_CMAC. The default value is AES_128_CMAC. You can configure a fallback key on the device to initiate a backup session if the primary session fails.
DCI MACsec Fallback Key String	Specify a Cisco Type 7 encrypted octet string that is used for establishing a fallback MACsec session. For AES_256_CMAC, the key string length must be 130 and for AES_128_CMAC, the key string length must be 66. If these values are not specified correctly, an error is displayed when you save the fabric.  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 [Create a VRF-Lite Inter-Fabric Link](#).

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

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

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, NDFC uses the fabric settings for configuring MACsec on the VRF-Lite IFC. When MACsec is configured on the link, NDFC 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, NDFC uses preshared keys provided by the user instead of using the QKD server to generate the keys. If you disable the Enable QKD option, all the fields pertaining to QKD are grayed out.
Use Link MACsec Setting	Check this check box as the override option for using the link settings instead of using the fabric settings.
MACsec Cipher Suite	Choose one of the following MACsec cipher suites for the MACsec policy: • GCM-AES-128 • GCM-AES-256 • GCM-AES-XPB-128 • GCM-AES-XPB-256 The default value is GCM-AES-XPB-256.
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.

Disabling 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 [Connecting Two Fabrics with MACsec Using QKD](#).

The following operations will 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, NDFC 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 **Security** tab and unselect the

Enable MACsec option.

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

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 Nexus 9000 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. Cisco NDFC support 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 Cisco Nexus N9K-C9332C Switch, Cisco Nexus N9K-C9364C Switch, Cisco Nexus N9K-C9348GC-FXP Switch as also the Cisco Nexus 9000 Series Switches that ends with FX, and FX2 support vPC fabric peering.
- Cisco Nexus N9K-C93180YC-FX3S and N9K-C93108TC-FX3P platform switches support vPC fabric peering.
- Cisco Nexus 9300-EX, and 9300-FX/FXP/FX2/FX3/GX/GX2 platform switches support vPC Fabric Peering. Cisco Nexus 9200 and 9500 platform switches do not support vPC Fabric Peering. For more information, see *Guidelines and Limitations for vPC Fabric Peering* section in *Cisco Nexus 9000 Series NX-OS VXLAN Configuration Guide*.
- If you use other Cisco Nexus 9000 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 Nexus 9000 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 NDFC.
- 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 Nexus 9000 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 the section "Creating a VXLAN EVPN Fabric Using the Data Center VXLAN EVPN Template" in [Data Center VXLAN EVPN](#).

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 the section "Creating a VXLAN EVPN Fabric Using the Data Center VXLAN EVPN Template" in [Data Center VXLAN EVPN](#).

To view the vPC pairing window of a switch, from the fabric topology window, right-click the switch and choose **vPC Pairing**. The vPC pairing window 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 window.
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:

Creating a Virtual Peer Link

To create a virtual peer link from the Cisco NDFC Web UI, perform the following steps:

1. Choose **Manage > Fabrics**.

The **LAN Fabrics** window appears.

2. Choose a fabric with the **Data Center VXLAN EVPN** or **BGP Fabric** fabric templates.
3. On the **Topology** window, right-click a switch and choose **vPC Pairing** from the drop-down list.

The window to choose the peer appears.



Alternatively, you can also navigate to the **Fabric Overview** window. Choose a switch in the **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** window, choose **Recalculate & Deploy**.

The **Deploy Configuration** window appears.

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

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

9. View the vPC link details in the pending configuration and side-by-side configuration.
10. Close the window.
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 window. For more information, see *Guidelines and Limitations for vPC Fabric Peering* and *Migrating from vPC to vPC Fabric Peering* sections in *Cisco Nexus 9000 Series NX-OS VXLAN Configuration Guide*.

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

Converting 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 Nexus N9K-C9332C Switch, Cisco Nexus N9K-C9364C Switch, and Cisco Nexus N9K-C9348GC-FXP Switch.
 - Cisco Nexus 9000 Series Switches that ends with FX, FX2, and FX2-Z.
 - Cisco Nexus 9300-EX, and 9300-FX/FXP/FX2/FX3/GX/GX2 platform switches. For more information, see *Guidelines and Limitations for vPC Fabric Peering* section in *Cisco Nexus 9000 Series NX-OS VXLAN Configuration Guide*.

To convert a physical peer link to a virtual peer link from the Cisco NDFC Web UI, perform the following steps:

1. Choose **LAN > Fabrics**.

The **LAN Fabrics** window appears.

2. Choose a fabric with the **Data Center VXLAN EVPN** or **BGP Fabric** fabric templates.
3. On the **Topology** window, right-click the switch that is connected using the physical peer link and choose **vPC Pairing** from the drop-down list.

The window to choose the peer appears.



Alternatively, you can also navigate to the **Fabric Overview** window. Choose a switch in the **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** window, choose **Recalculate & Deploy**.

The **Deploy Configuration** window appears.

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

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

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

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.

Converting 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.

To convert a virtual peer link to a physical peer link from the Cisco NDFC Web UI, perform the following steps:

1. Choose **LAN > Fabrics**.

The **LAN Fabrics** window appears.

2. Choose a fabric with the **Data Center VXLAN EVPN** or **BGP Fabric** fabric templates.
3. On the **Topology** window, right-click the switch that is connected through a virtual peer link and choose **vPC Pairing** from the drop-down list.

The window to choose the peer appears.



Alternatively, you can also navigate to the **Fabric Overview** window. Choose a switch in the **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** window, choose **Recalculate & Deploy**.

The **Deploy Configuration** window appears.

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

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

8. View the vPC peer link details in the pending configuration and the side-by-side configuration.
9. Close the window.
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 window.

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

Deploying Fabric Underlay eBGP Policies

To deploy fabric underlay eBGP policy, you must manually add the **leaf_bgp_asn** policy on each leaf switch to specify the BGP AS number used on the switch. Implementing the **Recalculate & Deploy** operation afterwards will generate eBGP peering over the physical interface between the leaf and spine switches to exchange underlay reachability information. If **Same-Tier-AS mode** is used, you can deploy the **leaf_bgp_asn** policy on all leafs at the same time as they share the same BGP ASN.

In a Multi-AS fabric, add the the **leaf_bgp_asn** policy on each leaf node and the fabric. In a vPC switch pair, they share the same AS number.

To add a policy to the required switch, see the section "Adding a Policy" in [About Fabric Overview for LAN Operational Mode Setups](#).

Deploying Fabric Overlay eBGP Policies

You must manually add the eBGP overlay policy for overlay peering. NDFC provides the built-in eBGP leaf and spine overlay peering policy templates that you must manually add to the eBGP leaf and spine switches to form the EVPN overlay peering.

Deploying Spine Switch Overlay Policies

Add the **ebgp_overlay_spine_all_neighbor** policy on the spine or super spine switches. This policy can be deployed on all the spine switches at once, since they share the same field values. If the network contains spine switches and super spine switches, you must deploy the policy only on the super spine switches.

1. Navigate to **Manage > Fabrics** and double-click on the fabric.

The **Fabric Overview** window appears.

2. On the **Policies** tab, choose **Actions > Add Policy**.
3. Select all the spine switches to which you want to add the **ebgp_overlay_spine_all_neighbor** policy and click **Next**.

The **Create Policy** window appears.

4. Click **Choose Template** and select the **ebgp_overlay_spine_all_neighbor** policy.
5. Enter the necessary field values for the following fields, as required and click **Save**.

Field	Description
Leaf IP List	Specifies the IPv6 or the IPv4 address of the connected leaf switch routing loopback interface. If you have enabled IPv6 underlay, ensure you enter the IPv6 address in this field.
Leaf BGP ASN	The BGP AS numbers of the leaf switches.
BGP Update-Source Interface	This is the source interface for BGP updates. Underlay routing loopback (loopback0) is used by default.

1. At the top-right of the **Fabric Overview** window, choose **Actions > Recalculate and Deploy**.
2. After the configuration deployment is complete, click **Close**.

You can use the **Edit Policy** option to edit the policy and click **Push Configuration** to deploy the configuration.

Deploying Leaf Switch Overlay Policies

Add the **ebgp_overlay_leaf_all_neighbor** policy on all the leaf switches, to establish eBGP overlay peering towards the spine switch. This policy can be deployed on all leaf switches at once, since they share the same field values.

1. Navigate to **Manage > Fabrics** and double-click on the fabric.

The **Fabric Overview** window appears.

2. On the **Policies** tab, choose **Actions > Add Policy**.
3. Select all the spine switches to which you want to add the **ebgp_overlay_leaf_all_neighbor** policy and click **Next**.

The **Create Policy** window appears.

4. Click **Choose Template** and select the **ebgp_overlay_leaf_all_neighbor** policy.
5. Enter the necessary field values for the following fields, as required and click **Save**.

Field	Description
Spine/Super Spine IPv4/IPv6 List	Specifies the IPv4 or IPv6 addresses of the routing loopback interfaces of spine or super spine switches for BGP peering. If you have enabled IPv6 underlay, enter the IPv6 address in this field. If the fabric has any super spine or border super spine, provide the IP addresses of the super spines or border super spines.
BGP Update-Source Interface	This is the source interface for BGP updates. Underlay routing loopback (loopback0) is used by default.

6. At the top-right of the **Fabric Overview** window, choose **Actions > Recalculate and Deploy**.
7. After the configuration deployment is complete, click **Close**.

You can use the **Edit Policy** option to edit the policy and click **Push Configuration** to deploy the configuration.

Adding a Super Spine Switch to an Existing VXLAN BGP EVPN Fabric

If your fabric contains both spine and super spine switches, you must reconfigure the fabric to use the super spine for deploying the overlay between the leaf and the border devices. This topic describes steps to integrate a super spine switch to your existing fabric which has a leaf switch and a spine switch with an overlay between them.

1. To add the **ebgp_overlay_spine_all_neighbor** policy to super spine switches that are newly added to an existing VXLAN BGP EVPN fabric:
 - a. Navigate to the **Fabric Overview** window for your fabric and click the **Policies** tab.
 - b. Select the super spine switches to which you want to add the **ebgp_overlay_spine_all_neighbor** policy and click **Next**.

The **Create Policy** window appears.

- c. Click **Choose Template** and select the **ebgp_overlay_spine_all_neighbor** policy.
- d. Enter the IPv4 or IPv6 addresses of the leaf switches in the **Leaf IP List** field.
- e. Enter the AS numbers for the leaf switches in the **Leaf BGP ASN** field and click **Save**.

2. To modify the existing **ebgp_overlay_leaf_all_neighbor** policy on each leaf node:
 - a. Find the existing policy by filtering based on **ebgp_overlay_leaf_all_neighbor** template name.



Ensure you modify only one policy at a time.

- b. Select a policy and choose **Action > Edit Policy**.
 - c. Enter the IP addresses of the super spine routing loopback interfaces in the **Spine/Super Spine IP List** field and click **Save**.
3. Select the leaf and the super spine switches that you have added in step 2 and choose **Actions > Deploy**.
4. On the **Links** tab, click on Protocol View and verify that eBGP peering between the super spine and leaf switches are established.
5. Remove the existing overlay between the leaf and the spine switches as follows:
 - a. On the spine switch, select the **ebgp_overlay_spine_all_neighbor** policy and choose **Actions > Delete Policy**.
 - b. On the leaf switch, select the **ebgp_overlay_leaf_all_neighbor** policy and choose **Actions > Edit Policy**.
 - c. Remove the IP address of the spine switch in the **Spine/Super Spine IPv4/IPv6 List** field and click **Save**.
6. To deploy the updated configuration on spine and leaf switches , choose **Actions > Recalculate and Deploy** at the top-right of the **Fabric Overview** window.

or

Select the leaf and the spine switches and choose **Actions > Deploy**.

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)

© 2017–2025 Cisco Systems, Inc. All rights reserved.